

Negative effects of redundant targets

Alex L. White^{*1}, John Palmer², Genevieve Sanders¹, Jannat Hossain¹
& Zelda B. Zabinsky³

¹Department of Neuroscience & Behavior, Barnard College, Columbia University.

²Department of Psychology, University of Washington

³Department of Industrial & Systems Engineering, University of Washington

* Corresponding author: Alex L. White, alwhite@barnard.edu

76 Claremont Ave, New York, NY 10031

Manuscript word count: 24,716 (including a 7,289 word Appendix)

ABSTRACT

The visual system can encode many stimuli simultaneously, but there are limits to how well multiple objects can be identified in parallel. At the extreme, some objects might have to be identified serially. The redundant target paradigm is one tool for distinguishing specific parallel and serial models. It compares responses to displays containing one target versus displays containing two targets. The typical result is a positive redundant target effect: faster correct responses to two targets, as predicted by many parallel models. Here we generalize three standard models to account for response accuracy as well as speed. Surprisingly, two models predict a reversal of the redundant target effect (*slower* responses to two targets than to one target): the generalized standard serial model, and a specific form of a fixed-capacity parallel model. To test that prediction, we measured performance for three different judgments of written words: color detection, lexical decision, and semantic categorization. The color task yielded positive redundant target effects, which reject the standard serial model. The semantic task yielded consistently negative effects, which are consistent with either the standard serial model or some limited-capacity parallel models. Thus, redundant targets can have negative effects, and they demonstrate limits that impair simultaneous recognition of two words.

Key words: redundant targets; divided attention; serial and parallel processing; response time; visual word recognition.

PUBLIC SIGNIFICANCE STATEMENT

The visual environment often contains many relevant objects, from faces in a crowd to words on a page. Here we improve an experimental test of whether you can process more than one such stimulus at a time. The results reveal limits to how well you can recognize two written words simultaneously.

The study of perception has long been animated by the question of how we process multiple stimuli when they are presented simultaneously. This question is particularly important for situations that involve visual search (e.g., locating a friend in a crowd) and for complex tasks such as reading, when many relevant stimuli (e.g., words) are presented simultaneously. Depending on several sensory and cognitive factors, an observer might be able to recognize multiple stimuli simultaneously, each just as well as a single stimulus viewed alone, or they might be hindered by a processing capacity limit. There are gradations of limited-capacity parallel processing, due to finite cognitive resources being shared between stimuli that might compete or interfere with each other. At the extreme, multiple stimuli might be processed serially, one at a time.

One of the most straightforward tools for evaluating parallel processing capacity is the redundant target paradigm (van der Heijden et al., 1983). In this paper, we revisit the redundant target paradigm and develop generalizations of standard models that predict both response times and errors. The first two models are the most common from prior work on redundant targets: a standard unlimited-capacity parallel model and a standard serial model. We also develop two variants of a fixed-capacity parallel model, which we also generalized to include accuracy. They make many of the same assumptions as the other two models (as explained in more detail below), but they introduce a cost of processing two stimuli in parallel. They can make different predictions depending on the nature of the evidence accumulation process.

We then test these predictions with four experiments that use written words as the target stimuli. These experiments assess how well two English words can be recognized simultaneously. The experiments presented here differ from natural reading, but they characterize the fundamental capacity limits that readers cope with. The redundant target paradigm complements other approaches that measure dual-task performance or spatial

attention effects to investigate how well readers can recognize two words at exactly the same time (Johnson et al., 2022; White et al., 2018, 2020; White, Palmer, et al., 2019).

Fundamentals of redundant target effects

The redundant target paradigm grew out of a larger visual search literature to investigate whether observers can process multiple stimuli presented simultaneously at different visual field positions (van der Heijden, 1975). The observer's task is to view a display and report the presence or absence of stimuli that belong to a target category. Non-target stimuli are termed "distractors." On some trials, one target is presented alone or with one or more distractors. On other trials, multiple targets are presented simultaneously – that is, the display contains redundant targets. The typical finding is a positive *redundant target effect*: a speeding of correct response times for multiple targets compared a single target. Such a "redundancy gain" is often taken as evidence that the targets were identified in parallel.

Studies that have used the redundant target paradigm can be divided into two broad categories. Studies in the first category seek to distinguish between a parallel model and a serial model (van der Heijden, 1975). They compare response times between displays that consist of two (or more) targets, versus displays that contain one target and no other stimuli. Positive redundant target effects have often been used to reject the serial model in favor of the parallel model. (Exceptions in the literature are reviewed below). Studies in the second category seek to distinguish between two flavors of parallel models: those with separate activations caused by each stimulus, versus those with "coactivations" (Eriksen et al., 1989; Miller, 1982; Mordkoff & Yantis, 1991). To do so, response times are compared between displays that contain two targets and displays that contain one target and one distractor. These "mixed" trials are not useful for testing the serial model, because on a random half of trials, the serial process would start with the distractor before identifying the target. This causes slower responses on average for

mixed target-distractor trials than two-target trials. That is the same prediction as made by parallel models.

Therefore, we focus on the comparison between displays with a single target presented alone and displays with two targets. As shown in **Figure 1**, contrasting predictions for correct response time arise from a standard unlimited-capacity parallel model and a standard serial model. They are called “*standard*” models because of assumptions about the independence of the processes for each stimulus. In particular, they assume “selective influence” of each stimulus on a processing channel, which means that the specific features of one stimulus do not influence the features that are represented for another stimulus. Both standard models assume that search is *self-terminating*: the observer responds as soon as they detect a target (Van Zandt & Townsend, 1993). In the General Discussion we address how parallel and serial models can mimic each other when made more complicated (Algom et al., 2015). In this study, however, we focus on relatively simple models with plausible assumptions, which we develop further by accounting for errors as well as response times. To make progress we then reject specific models that cannot account for performance in word recognition tasks.

The standard, self-terminating, unlimited-capacity parallel model assumes that when two targets are present, they are independently processed in separate channels that race to produce the response. The completion time of each process is variable across trials. On two-target trials, the response time is determined by the faster of the two processes, so the observer is faster on average than when only one target is present (Raab, 1962). Thus, the standard, self-terminating parallel model predicts a positive redundant target effect: a speeding of correct responses. Note that the term “*unlimited-capacity*” as applied to this standard parallel model does not imply that an unlimited *number* of stimuli could be processed in parallel. Rather, it assumes that in the specific context of the experiment, two stimuli are processed independently just as well as a single stimulus is processed,

without a capacity limit. It is the common “race” model often discussed for interpreting redundant target effects.

The standard self-terminating serial model, in contrast, assumes that one stimulus is identified at a time (Townsend & Nozawa, 1995; van der Heijden, 1975). If two targets are presented simultaneously, and the first one to be processed is identified correctly, then the mean response time is the same as when only one target is present—the redundant target has no effect on performance. Previous descriptions of this serial model stop there; but as we explain in our new theory section below, a serial model that incorporates errors predicts a *slowing* of correct response times if the first target to be processed is misidentified, and search continues to process the second target correctly.

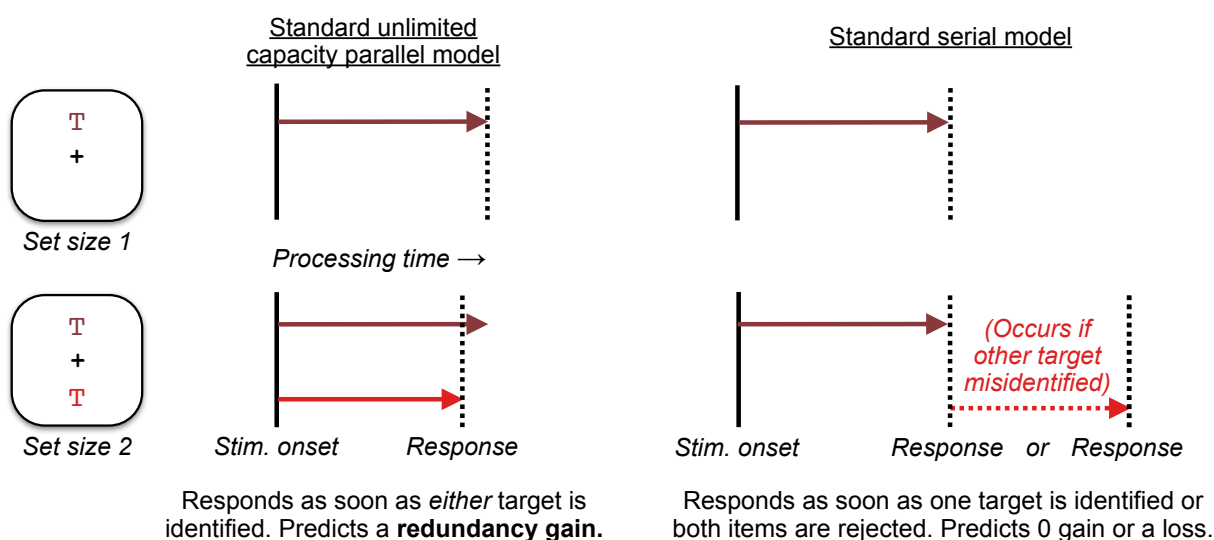


Figure 1: Diagram of a standard parallel model and a standard serial model for displays containing one target (set size 1) or two targets (set size 2). The parallel model predicts faster correct responses for set size 2 because the response is triggered by whichever process happens to finish sooner. The serial model predicts slower responses to two targets as on some trials one target is misidentified as a distractor before the other target is identified.

Positive redundant target effects have been used to reject the standard self-terminating serial model for processing simple visual features, such as detecting lights, discriminating colors, orientations, and motion directions (Corballis, 2002; Donkin et al., 2014; Egeth et al., 1989; Ridgway et al., 2008; Schwarz, 2006; Thornton & Gilden, 2001).

Redundant target effects have also been found with auditory stimuli (e.g. Schröter et al., 2007) and with bimodal stimuli (e.g. Gondan et al., 2010; Hershenson, 1962). One study reported no significant redundant target effect in a face recognition task (Fitousi, 2021).

Letters are an interesting case, being the building block of words. Several studies have used tasks that require the participant to distinguish one target letter from other letters. When the task uses a “go-no/go” procedure—the task is to press a button when a target is detected and otherwise make no response—there are positive redundant target effects (e.g. Grice & Reed, 1992; Mordkoff & Yantis, 1991; van der Heijden et al., 1983). That is also true when the observer makes a vocal choice response (van der Heijden, 1975). However, other studies have found no redundant target effect when the procedure is slightly different, such as requiring a button-press choice response on each trial (Fournier & Eriksen, 1990; Grice & Reed, 1992; van der Heijden et al., 1983). In some of the redundant target literature, variation in the task procedure (go-no/go or choice) was confounded with the inclusion or exclusion of mixed target-distractor trials. But two studies using letters as stimuli do show differences between go/no-go and choice tasks while holding the rest of the design constant (Van der Heijden, et al., 1983; Grice and Reed, 1992). Thus, for letters, there is reasonable evidence that positive redundant target effects are more robust for the go/no-go task procedure. One possible explanation is that letters are processed in parallel, producing a positive redundant target effect, but that effect can be masked by later decision- or response-selection processes when the response rule is more complicated. That is one reason we tested both types of procedures in the present study.

Redundant target effects for word recognition

We now turn the case explored in the experiments below: redundant target effects for written words. Such effects can reveal the extent to which higher-level semantic or linguistic information about two stimuli can be processed in parallel. Of the handful of studies that have taken this approach, most – but not all – conclude that two words can

be recognized in parallel. These studies have differed in four important respects: whether the redundant targets in a single trial are identical words; whether the single-target trials also contain a 'filler' stimulus; what the task is (semantic categorization vs. lexical decision), and how the subject responds (go/no-go vs. choice).

We first summarize experiments that reported a positive redundant target effect for word recognition. Shepherdson & Miller (2014) used a semantic categorization task. Words were presented to the left and/or right of fixation in some experiments and above and/or below fixation in others. On each trial, the subject had to make a yes/no response to report the presence of a word belonging to a target semantic category (animals). Importantly, the "single-target" trials also contained a "filler" stimulus that was a pronounceable pseudo-word (neither a target, nor a distractor of the other semantic category). The key result was that responses were faster in the redundant target condition compared to this modified target-filler baseline.

Interpreting this experiment is difficult. The critical question is what the standard self-terminating serial model predicts in terms of the response time difference between these two conditions. It predicts 0 difference only if we make an additional assumption about the trials with a single target paired with a pseudoword filler: the target stimulus is always processed first and the pseudoword is never processed at all (at least on correct trials). That seems unlikely, as the observer would have to first determine which stimulus is a real word before beginning to process it. Alternatively, if we assume that the serial process sometimes begins with the filler stimulus before moving on to the target word, then the serial model predicts slower responses on these target-filler trials than on two-target trials. The serial model thus predicts the same thing as most parallel models (a positive redundant target effect). Thus, we conclude that Shepherdson & Miller (2014) did not strictly test the standard self-terminating serial model that we wish to test. Instead that study is part of the broader set of experiments that use a mixed-pair baseline to test coactivation models.

We must go further back in time to find studies that did test the standard serial model for word recognition by comparing displays containing two words to display containing one word and nothing else. In the four experiments reported by (Mullin & Egeth, 1989), words were presented above and/or below fixation, and the participant's instruction was to make a go/no-go response to the presence of any target. Trials contained 1 distractor alone, 1 target alone, 2 distractors alone, or 2 targets alone. In two experiments with a semantic categorization task, the redundant target effect did not significantly differ from 0. This is unlike the many other experiments with tasks based on simpler stimulus features. In additional experiments they used a lexical decision task: targets were real English words and distractors were pseudowords. When the redundant targets were two identical copies of the same word, there was a significantly positive redundant target, consistent with parallel processing. A similar result was reported by (Egeth et al., 1989). Other studies found consistent positive effects for lexical decision with identical words present to the left and right of fixation: Hasbrooke & Chiarello (1998) and Mohr and colleagues (Mohr et al., 1994, 1996).

However, a redundant target effect for identical words might be explained by facilitation at a sub-lexical level (Abrams & Greenwald, 2000; Grainger et al., 2014)¹. Thus, in we focus on Mullin and Egeth (1989)'s second lexical decision experiment, in which the two words on redundant target trials were different from each other. The redundant target effect was then significantly *negative*, meaning that response times were slowed by the addition of a second target.

In summary, most redundant target studies of word recognition found a positive effect, although there are two examples of either zero effect or a negative effect. There was evidence for positive effects when the redundant targets were identical words but not when they were two different words. In the study below we focus on parallel

¹ Nonetheless, presenting copies of the same word across the field provides a redundancy gain that might help patients with macular degeneration to read, with rapid serial visual presentation (Snell et al., 2022).

processing of two non-identical targets. This is motivated by a long-term goal of understanding processing capacity limits that affect reading, when neighboring words are not identical. The literature reviewed above also indicates that the direction of the redundant target effect might depend on the subject's task (lexical decision or semantic categorization), on the mode of response (go/no-go or choice), and on whether the experiment includes mixed trials in which a target is paired with a distractor. In the new experiments reported below, we investigate all those factors, and compare lexical and semantic tasks to a font color task. In doing so we also test new models of parallel and serial processing that consider both response time and accuracy.

Response time and accuracy in redundant target effects

Most of the redundant target studies reviewed above focused on only correct response times and use tasks in which accuracy is near ceiling. Other studies have focused on accuracy, for instance in the context of spatial summation (Robson & Graham, 1981; Verghese & Stone, 1995). One important study compared redundant target effects on accuracy and response time (Mordkoff & Egeth, 1993). This work has shown that typical parallel models predict positive redundant target effects for accuracy as well as response time.

As shown in the following section on our new theory, serial processing of the individual stimuli can lead to a *negative* effect of redundant targets on response time. This novel result hinges on the possibility of errors: if the first target to be processed is misidentified as a distractor, then search continues, and the second target may be correctly identified. Those correct responses increase the mean response time for two-target displays compared to correct responses to single target displays. Thus, our new theory explicitly considers the accuracy of each stimulus recognition process when predicting response times.

New generalizations of standard models generate novel predictions

The new theory generalizes previous models of pure response time by adding the possibility of errors (misclassifying targets as distractors or vice versa). We focus on models that for decades have been central to the literature on visual search and redundant targets: the standard unlimited-capacity parallel model and the standard serial model. We then also develop variants of fixed-capacity parallel models, which are novel in the study of redundant target effects. The **Appendix** contains full mathematical descriptions of these three classes of models. Here we describe them in intuitive terms and emphasize the qualitative redundant target effects that they each predict.

Our goal in building this new theory is to compare the *qualitative* predictions of the various models: whether they predict positive, negative, or zero redundant target effects on response time and accuracy. Our goal is not to quantitatively fit models to our data; that is a different endeavor that we leave for the future (see also Cox & Criss, 2019). For now, it is sufficient to generate qualitative predictions that allow experimental data to rule out some models.

Figure 2 illustrates the typical range of predicted redundant target effects for each class of model that incorporates errors. The standard serial model predicts negative effects. The standard unlimited-capacity parallel model predicts positive effects. The standard fixed-capacity parallel model can predict either negative or positive effects. These results of the new theory are described in more detail in the following paragraphs.

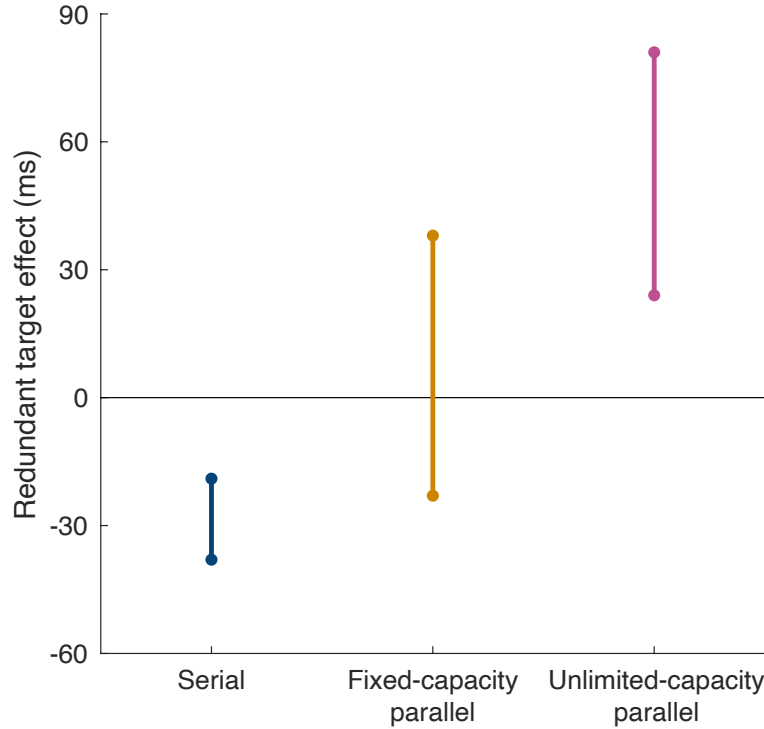


Figure 2: Model predictions. For each type of model, we plot the range of predicted redundant target effects on correct response times. For all models we assume that for single-target trials, errors occur on 5% of trials and the mean correct response time is 600 ms. Each model's range encompasses the smallest and largest effects that we generated with the parameters described at the end of the Appendix.

Predictions of Our Serial Model

We start with the standard self-terminating serial model developed for response time (Townsend & Nozawa, 1995). Like others of its type, it assumes discrete component processes for each stimulus. In addition, it allows time for residual processes before a response is made that do not depend on the stimulus. This model has been called standard because it includes a number of independence properties. For instance, the component processing time for one stimulus is independent of the component processing time for the other stimulus, and of the presence or absence of the other stimulus and the other stimulus's features (see Appendix). We add to this model the possibility of an error

(i.e., a stimulus is judged incorrectly) and additional independence assumptions concerning the errors.

The most important result is that this model predicts a *negative* redundant target effect as long as accuracy is below 100% correct. This is unlike the corresponding pure response time model that predicts no effect of redundant targets. This different result arises from errors in processing one of the two targets: if the first target to be processed is mis-identified as a distractor, processing continues for the second stimulus. If this second target is correctly identified, then this correct response time is included in the analysis with the other trials in which the first target was processed correctly. Thus, the correct response times for the redundant target condition are a mixture of two cases: those trials in which the first target was correctly identified and the response was made quickly, and those trials in which the first target was not correctly identified and processing continued to the second target. This new result makes the serial model with errors more distinctive from the parallel models, in terms of correct response times. The predictions for errors are discussed below.

Quantitatively, this model's predicted redundant target effect is given by Equation 7 in the Appendix, which is repeated here:

$$\mu_{t,correct} - \mu_{tt,correct} = -\frac{(1 - p_t)}{(2 - p_t)} E[\mathbf{D}_{t,incorrect}].$$

The redundant target effect is the difference between the mean correct response time for a single target ($\mu_{t,correct}$) and the mean correct response time for two targets ($\mu_{tt,correct}$). This effect depends on only two factors, the probability correct on single-target trials, p_t , and the mean component processing time for a single target when the participant makes an error (a 'miss' response), $E[\mathbf{D}_{t,incorrect}]$.

Crucially, the serial model predicts a negative redundant target effect whenever there are errors (i.e., whenever some targets are misidentified as distractors). An

illustration of this prediction is in **Figure 2** (leftmost bar). For this illustration, we assume that for single-target trials, accuracy p_t is 0.95 (5% errors) and the mean correct response time is 600 ms. As described in the Appendix, we assume a residual time of 200 ms, which yields a component processing time of 400 ms. The upper end of the range is predicted with the assumption that the mean component processing time for an error is equal to the mean component processing time for a correct response. The lower end of the range is predicted with the assumption that the mean component processing time for an error is twice that for a correct response.

Predictions of Our Unlimited-Capacity, Parallel Model

Our second result concerns the standard, self-terminating, unlimited-capacity parallel model with errors. It assumes that two stimuli are processed in parallel and independently. There are two forms of independence for this model: first, “selective influence”: each stimulus is processed in a spatially selective channel that is not influenced by the features of other stimuli. Second, the speed and accuracy of each stimulus process is not affected by the absence or presence of any other stimuli that are simultaneously processed.

As previously developed, this model without errors predicts a positive redundant target effect. We show that this positive effect generalizes even when errors occur. Errors can reduce the size of the effect, but it always remains positive. Thus, there remains a sharp contrast in the predictions for this parallel model and the standard self-terminating serial model.

The predictions of this parallel model are given by Equation 17 in the Appendix:

$$\mu_{t,correct} - \mu_{tt,correct} = \left(\frac{1}{2 - p_t}\right) E[\mathbf{D}_{t,correct}] - \left(\frac{p_t}{2 - p_t}\right) E[\min\{\mathbf{D}_{t_1,correct}, \mathbf{D}_{t_2,correct}\}]$$

Here, p_t is the probability correct on single-target trials, $E[\mathbf{D}_{t,correct}]$ is the mean correct component processing time for a single target, and $E[\min\{\mathbf{D}_{t_1,correct}, \mathbf{D}_{t_2,correct}\}]$ is the mean of

the minimum of the two component processing times when two targets are presented and judged correctly.

The predicted effect is always positive. This is primarily because the model's response to two targets is driven by whichever of the two stimulus processes finishes first, hence the “min” function in the equation above. On average this is faster than the response to a single target. The equation also makes clear why the predicted effect is always positive: it is the difference between two products, and the first is always larger. This must be the case, as is clear when examining each part of the two products separately. First:

$$\left(\frac{1}{2 - p_t}\right) \geq \left(\frac{p_t}{2 - p_t}\right)$$

because $0 \leq p_t \leq 1$. Second:

$$E[\mathbf{D}_{t,correct}] > E[\min\{\mathbf{D}_{t_1,correct}, \mathbf{D}_{t_2,correct}\}]$$

because the mean difference between two identically distributed (non-negative) variables is always less than the mean of one of those variables alone.

An illustration of this prediction is the rightmost bar in Figure 2. For this illustration, we used predictions of two specific models described in the Appendix. They differ in the nature of the stochastic evidence accumulation process: how activation of each stimulus channel triggers a decision to respond. The upper point is for a diffusion model with parameters that generate large redundant target effects. The lower point is for a linear ballistic accumulator model with parameters that generate relatively small redundant target effects. While not strict limits, these model outputs illustrate the range of predictions from the standard, self-terminating, unlimited-capacity parallel model. They are always positive.

Predictions of Our Fixed-Capacity, Parallel Models

The parallel model that can most easily mimic a serial model is one that has limited capacity. The limited capacity slows processing when there are two stimuli and thus reduces and possibly eliminates the redundant target effect. Unfortunately, this model is so general that it does not make very specific predictions. Here, we consider a special case of the limited-capacity parallel model: the fixed-capacity parallel model. The idea of fixed capacity is that a set of parallel processors extract the same total amount of information from multiple stimuli as they do from a single stimulus. Thus, splitting a fixed set of ‘resources’ between multiple stimuli introduces a cost. Most of the prior work with this model has been in the domain of accuracy (Scharff et al., 2011; Shaw, 1980; White et al., 2018). Here we also introduce a slowing to processing *time* when a fixed amount of resources are divided between two targets, as compared to when only one target is present.

We investigated two special cases of self-terminating fixed-capacity parallel models in which we assume a particular stochastic evidence accumulation process for each stimulus being processed. Predictions of these two special cases define the range of redundant target effects plotted in Figure 2 (middle bar). First, with a diffusion process of sensory evidence accumulation (Palmer et al., 2005), the model yields positive redundant target effects on correct response times, for all relevant parameter values (as well as a positive effect on accuracy). That is because each stimulus component process has variable completion times, but when two targets are present, two processes race to trigger the response (as for the unlimited-capacity parallel model). That leads to a statistical benefit on average, which can outweigh the slowing caused by fixed capacity. This prediction defines the upper end of the range of effects predicted by the fixed-capacity parallel model in Figure 2.

However, with a linear ballistic accumulator (LBA) process (Brown & Heathcote, 2008), the fixed-capacity, parallel model can predict a *negative* redundant target effect on

correct response time (a slowing), despite a positive effect for accuracy. That is because the component processing times are less variable in the LBA model than the diffusion model. Thus, there is less of a benefit from the “racing” of two parallel processes, and the slowing due to the capacity limits reveals itself as a negative redundant target effect. This is illustrated in Figure 2 by the lower end of the range of predicted effects for the fixed-capacity model. The key point is that among many parallel models that generate positive response time effects of redundant targets, there are parallel models with fixed capacity that yield the opposite result.

Thus, there is an asymmetry in using the redundant target paradigm to test the serial model and fixed-capacity parallel models. All the models we consider here are assumed to be standard, self-terminating models. A positive redundant target effect rejects this serial model, but a negative redundant target effect does not reject all fixed-capacity parallel models. Thus, the redundant target paradigm is an excellent test for rejecting both the serial model and the unlimited-capacity parallel model, but it cannot distinguish the serial model from some fixed-capacity parallel models. Nevertheless, it is relevant to distinguishing specific serial and parallel models. Indeed, many experiments have used this test to rule out the standard self-terminating serial model (e.g., van der Heijden, 1975), and most redundant target studies of word recognition also claim to have rejected the serial model (as reviewed above).

Predictions About Errors

All of the models described above predict a positive redundant target effect on accuracy (that is, fewer errors on trials with 2 targets than on trials with 1 target). This is not a surprise for typical parallel models that have been investigated in summation experiments of accuracy alone (Graham et al., 1978). What is new is that this result also occurs for our serial model, even though that model predicts slower response times for two targets. The reason is that when two targets are present and processed sequentially,

there are two chances to correctly detect target presence, so accuracy increases compared to when only 1 target is present – even though doing so takes more time on average. While this result for errors does not distinguish between the serial and parallel models, it introduces a result that is specific to errors and that is not accounted for by pure response time models.

Predictions about response time variability

These models also make differing predictions about the distributions of correct response times. The Appendix summarizes these predictions within the sections for each model. The final section of the Appendix also includes simulations of response time distributions. The standard serial model assumes that on some trials, one target is misidentified and then the second target is processed correctly. Those trials add long response times to the distribution. However, our simulations show that the standard serial model does not clearly predict a *bimodal* distribution of correct response times on two-target trials. Instead, it predicts a distribution with a longer rightwards tail and a larger standard deviation than the distribution for single target trials. The redundant target effects on response time variability are small relative to the observed variability within each condition. The unlimited-capacity parallel model predicts the opposite: decreased variability for redundant target trials. Different flavors of the fixed-capacity parallel model can predict either an increase or decrease in the variability of response times.

In the results below, we do not analyze response time distributions or variability, because the experiments were not designed to detect such effects. Those effects are predicted to be small relative to the variability within each condition and difficult to detect empirically. Moreover, the model predictions for increases or decreases in variability mirror the predictions for the response time means, which our experiments were optimized to measure.

Overview of the experiments

We conducted four experiments that differed in three factors, which are meant to span the range of what is typical in redundant target studies. The first factor was how participants responded to the stimuli. “Go/No-Go” is a procedure that requires the participant to press a button if they see a target and to make no response if they see no target. This is the most common procedure in the redundant target literature. “Choice” is a procedure that requires the participant to press one of two buttons to categorize each stimulus display. This is the most common procedure in the larger visual search literature. As discussed above, some prior research suggests that a go/no-go procedure is more sensitive for detecting redundant target effects (Grice & Reed, 1992; van der Heijden et al., 1983). Previous studies about word recognition have used a mix of go-no/go and choice procedures, so we used both in different experiments.

The second factor we manipulated was whether the words presented on two-word trials were “correlated.” In the “correlated” design, the two words were either both targets or both distractors. In the “uncorrelated” design, there also were trials in which one target was paired with one distractor. The correlated design maximizes the fraction of trials that test the standard serial model (one target alone vs two targets). The uncorrelated design requires more trials but is more typical in visual search generally. The inclusion of mixed trials might affect the participant’s strategy and encourage them to process both stimuli, thus we use it in Experiments 3 and 4 to compare to the correlated design. As elaborated in the General Discussion, the uncorrelated and correlated designs have different contingencies (i.e., the probability of a target at one location given what is at the other location) that might change the participant’s strategy (Mordkoff & Yantis, 1991). But within each experiment, the contingencies were the same for the three tasks (lexical, semantic, and color).

The third factor we manipulated was the position of the words. In Experiment 1-3, we presented words directly above and below fixation in order to match the design of Mullin & Egeth (1989). This allows both words to be close to fixation and easily legible. In Experiment 4, we sought to make the displays more like natural English reading by arranging the words horizontally, just to the left and right of fixation, with a single letter space between them.

As noted above, when two words were present, they were always two different words. This differs from some previous redundant target experiments that used identical words on two-target trials (Hasbrooke & Chiarello, 1998; Mohr et al., 1994; Mullin & Egeth, 1989) and might produce effects due to sub-lexical facilitation. The one published case of a negative redundant target effect was in a lexical decision task with non-identical redundant targets (Mullin & Egeth, 1989)

In each experiment, we also measured performance in three different tasks (color detection, lexical decision, and semantic categorization), all with similar stimuli. The *color task* required participants to judge a low-level visual feature of the words, and served as a positive control condition for which we expected positive redundant target effects. The *lexical decision task* requires the subject to distinguish real English word targets from pseudoword distractors. The *semantic categorization task* requires categorizing words either as targets that belong to a category of “living things” and distractors that belong to a category of “non-living things”. The semantic and lexical tasks might tap into different levels of linguistic processing, and have both been used in prior redundant target studies (Egeth et al., 1989; Mullin & Egeth, 1989). In sum, within each of our four experiments, we carry out a side-by-side comparison of redundant target effects that arise in three tasks using similar stimuli. The tasks differ in which they require low-level feature detection, lexical access, or semantic categorization. Altogether, this study includes over 255,000 trials of data from a total of 326 participants.

General methods

Participants: Participants were recruited from around the world using Prolific (www.prolific.co, accessed August 2021-April 2025). Participants gave informed consent in accordance with the Declaration of Helsinki and Barnard College's Institutional Review Board. All participants indicated being fluent speakers who learned English as their first language, with no literacy difficulties, and normal or corrected-to-normal vision. For each task, we aimed to recruit an independent sample of 28 participants, half male and half female.

The sample size was chosen on the basis of a power analysis of an independent pilot data set with the same design as the color task for Experiment 2 as described below. (Experiment 2 was actually conducted first). 11 volunteers from the Barnard College and Columbia University community participated, with the same inclusion criteria as in the main experiment. The mean redundant target effect was 29 ms, with a standard deviation (SD) of 17 ms. To determine the sample size for each task in the experiments, we calculated the minimum number of participants required for 95% power to reject the null hypothesis with a mean effect of 14 ms (half as large as in the pilot set) and with the same standard deviation. This minimum was 24 participants. In the first experiment we conducted (Experiment 2's color task), an error in the online recruiting platform required us to recruit an additional 4 participants to reach a roughly equal number of males and females.

After conducting Experiment 2, we ran another power analysis based on the observed the redundant target effect for the *semantic* task (mean = -17.5, SD = 24 ms). This demonstrated that N=28 suffices for 95% power to detect such a negative effect (for 90% power, N=23). Thus, we sought a sample size of 28 for all the experiments.

Table 1 indicates the number of subjects and exclusions for all experiments in the study. Across all experiments in this study, our criteria for excluding participants were: overall d' less than 0.5, or more than 10% of trials excluded for response times being too

fast or too slow. In sum, 20 of 346 participants were excluded from the analysis. d' is a theoretically bias-free measure of accuracy that combines the hit rate (proportion correct on target-present trials) and false alarm rate (proportion incorrect on target-absent trials). An overall d' of 0.5 corresponds to roughly 60% correct (assuming a neutral criterion). 15 participants were excluded because $d < 0.5$. In all experiments, individual trials were also excluded if the response time was too fast (<250 ms from stimulus onset) or too slow (>3000 ms). The mean percentages of trials excluded for those reasons (amongst included participants) are listed in Table 1. Any participant for whom more than 10% of trials met those exclusion criteria was excluded entirely from the data set. In total, 5 participants were excluded for this reason alone.

Experiment	Task	N tested	N excluded	N included	N Female	Mean age [min max]	RTs too fast	RTs too slow
1: Go/No-Go, correlated	Color	30	1	29	12	32 [19 50]	0.27%	0.00%
	Lexical	28	0	28	12	32 [20 47]	0.19%	0.00%
	Semantic	28	0	28	14	30 [20 47]	0.06%	0.00%
2: Choice, correlated	Color	29	0	29	16	27 [18 50]	0.10%	0.09%
	Lexical	28	0	28	14	29 [19 48]	0.05%	0.42%
	Semantic	28	0	28	18	27 [19 50]	0.13%	0.41%
3: Choice, uncorrelated	Color	28	0	28	22	19 [18 24]	1.50%	0.27%
	Lexical	29	3	26	14	31 [20 48]	0.22%	0.40%
	Semantic	29	1	28	14	31 [20 48]	0.28%	0.49%
4: Left/Right Choice, uncorrelated	Color	30	8	22	10	32 [18 47]	1.63%	0.38%
	Lexical	29	3	26	9	31 [18 47]	0.63%	0.60%
	Semantic	30	4	26	12	31 [18 48]	0.82%	0.59%

Table 1: Number of participants (N) in each experiment and task, as well as the ages in years of the included subjects. The two rightmost columns are the mean percentages of trials that were excluded for response times (RTs) being too fast (<250 ms) or too slow (>3 s), within the included participants.

From Prolific we obtained self-reported data on race or ethnicity for 290 of the included participants. Of those, 9% were Asian, 22% were Black, 5% were more than one race, 61% were White, and 3% were “other.”

Transparency and Openness: Experiments were programmed using PsychoPy 3 (Pierce et al, 2019) and run online with Pavlovia.org. Data were analyzed in MATLAB 2022a (Mathworks, Inc), using the bayesFactor toolbox (<https://doi.org/10.5281/zenodo.4394422>). All data, analysis code, and stimulus materials have been made publicly available at the Open Science Framework and can be accessed at <https://osf.io/7kn9u/>. Above we report how we determined our sample size, all data exclusions, all manipulations, and all measures in the study. This study’s design and its analysis were not preregistered.

Experiment 1: Go/No-Go procedure with correlated stimuli

Methods

Stimuli: We created and presented stimuli with PsychoPy 3 (Peirce et al., 2019), run through the web browser using Pavlovia (<https://pavlovia.org/>). Each stimulus size and position were defined as a fraction of the height of the participant’s screen; thus, the dimensions in degrees of visual angle likely varied across participants. Participants were asked to sit with their head roughly 1 arm’s length from their screen. A central black fixation cross, 4.5% of screen height in width, was present throughout each trial except during feedback. The stimuli consisted of letter strings, of length between 4 and 6 letters. The word lists are described below and provided fully in the public data repository (<https://osf.io/7kn9u/>). They were drawn in Courier font, with the height of an “o” or “x” occupying 3.8% of the screen height. That is roughly 0.7 degrees visual angle on a typical laptop at arm’s length. The specific words (or pseudowords) and font color varied across tasks, as described below.

Trial sequence: An example trial is illustrated in **Figure 3A**. Each trial began with the fixation mark present for 750 ms. Then either 1 or 2 words were presented for 183 ms. There were two possible word positions, centered horizontally and either just above or just below the fixation mark. The distance from the fixation mark to the center of each word was 10% of the screen height (roughly 1.8 degrees visual angle on a typical laptop at arm's length). About two letter o's would fit stacked vertically in the empty space between the screen center and the words.

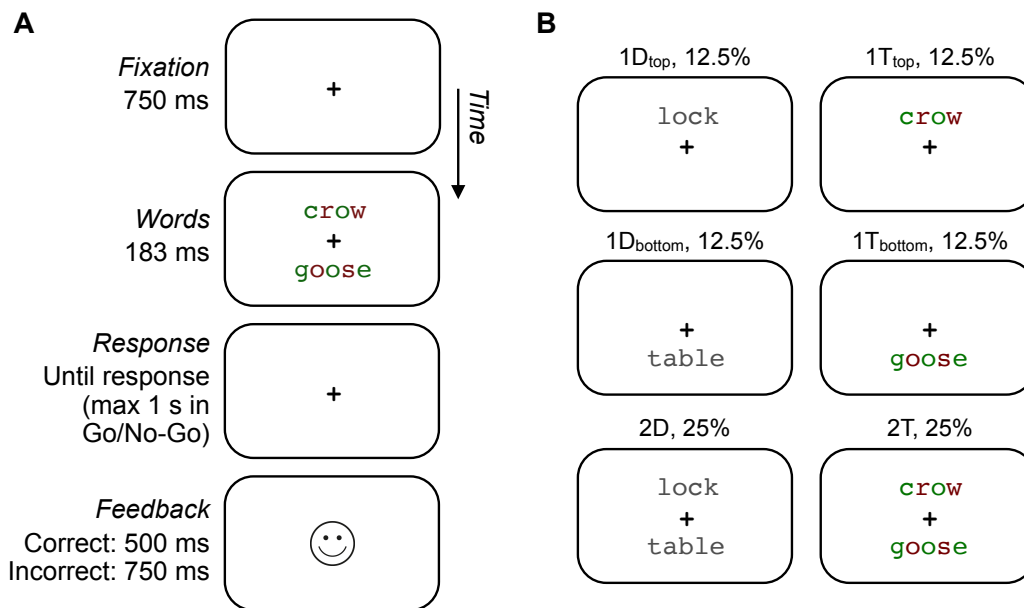


Figure 3: Stimuli and Design of Experiments 1 and 2. (A) Example trial sequence with two color targets. **(B)** Examples of the trial types for the color task. The text above each panel indicates the percentage of trials that were of that condition. “D” = distractor, “T” = target.

The trials were evenly distributed between these 4 conditions: 1 target, 1 distractor, 2 targets, and 2 distractors. **Figure 3B** illustrates examples of each trial type, and Table 2 lists the proportion of trials assigned to each combination of stimuli at the top and bottom locations. Each location could have no word (‘None’), a distractor word, or a target word. When just 1 word was present, it was equally likely to be in the top or bottom location. There were never any mixed pairs of 1 target and 1 distractor (unlike in Experiments 3

and 4). Thus, in this experiment, the words in each display were “correlated,” meaning that when there were 2 words present, they were either both distractors or both targets.

After the words disappeared, the participant was free to respond. In these go/no-go tasks, the participant was instructed to press the spacebar as soon as they detected a target, and to do nothing if they saw no targets. Up to 1 second was allowed for a response. After the response interval elapsed or was ended by a keypress, feedback was given: the fixation cross was replaced with a smiley face for 500 ms if the response was correct, or a neutral face for 750 ms if the response was incorrect. Then, the fixation cross reappeared and another trial began (except when it came time for a break between blocks, see below).

		Bottom word		
		<i>None</i>	<i>Distractor</i>	<i>Target</i>
Top word	<i>None</i>	N/A	0.125	0.125
	<i>Distractor</i>	0.125	0.25	0
	<i>Target</i>	0.125	0	0.25

Table 2: The probability of stimulus pairings at the top and bottom locations in Experiments 1 and 2. The word at each location was either absent, a distractor, or a target. The green shading highlights conditions when 2 words were present. In this design, the two words were perfectly correlated, meaning that they were either both targets or both distractors.

Procedure: Once they accessed the experiment in Pavlovia, participants read a consent form and indicated their acceptance by pressing a key to continue. The program advanced through four pages of instructions with example stimuli. Then the participant conducted practice trials, which continued for at least 50 trials until the participant had responded correctly to 36 of the most recent 40 trials. Having completed that, they began the main experimental trials which came in 10 blocks of 60 trials each. Before starting the

first block, participants were reminded to keep their head 1 arm's length from the screen, maintain central fixation, and to respond as quickly as possible without making unnecessary errors. Between each block they were given written feedback about their percent accuracy (P) and the opportunity to rest. If P for the most recent block was at least 96%, the feedback said, "Very nice! You got $\{P\}$ % correct. In the next block, try to go a bit faster, while still getting at least 90% correct." If $P \leq 72\%$ correct, the feedback said, "Good job. You got $P\%$ correct. In the next block, try to get above 90% correct." Otherwise, the feedback simply said, "You're doing great!" Participants completed the whole experiment in roughly 30 minutes, on average.

Color detection task: Each word was either drawn in all dark gray letters (RGB 79, 79, 79 out of 255) or its letters alternated between dark red (RGB 115, 18, 18) and dark green (RGB 17,102,15). These colors were 85% saturated, relative to the maximum possible on the screen. Targets were defined as words written in colored letters; distractors were words written in gray letters. The words were drawn from the same set as in the semantic categorization task (see below).

Lexical decision task: All the letters were dark gray (RGB 79, 79, 79). There was a total of 246 items in the stimulus set, half real English words and half pronounceable pseudowords. Within both categories, 33 had 4 letters, 46 had 5 letters, and 44 had 6 letters. The real words were all nouns that were also used in the color & semantic tasks, with mean lexical frequency 16.4 occurrences per million (ranging 0.3-391). The pseudowords were generated using MCWord (Medler & Binder, 2005) to have trigram statistics (the probability of any sequence of three letters) matched to real words. Across the 600 trials in the experiment, each word was repeated on average 3.7 times. We took care to match the mean lexical frequency and word lengths across trials with 1 real word and trials with 2 real words.

Semantic categorization task: All the letters were dark gray as in the lexical task. There was a total of 246 English nouns in the stimulus set, half of which referred to living

things and half to non-living things. Within the living category there were 39 4-letter words, 42 five-letter words, and 42 six-letter words. They referred to animals (e.g., “bird,” “turtle,” “woman”) and plants (e.g., “fern,” “orchid”; and one was “fungus”). The non-living category had 37 4-letter words, 42 five-letter words, and 44 six-letter words. They referred to common household items (e.g. “towel”), pieces of clothing (“e.g. “shoe”), and types of buildings (e.g., “cabin”), as well as natural non-living things (e.g., “snow”). The distributions of lexical frequencies in the living and non-living categories were highly overlapping, with means 19.7 and 14.5, respectively. Each word was repeated on average 3.7 times within the experiment.

Analysis: We computed two measures of performance in each condition: the mean response time on correct trials, and the percent of trials with incorrect responses (errors). In most cases we focus on trials with targets, because only those measures test our models that assume self-terminating search for targets. For both measures, we compared the means on trials with two targets to trials with one target using paired t-tests. All t-test p-values were corrected for false discovery rate across the 12 tests done for each measure in the entire study (Benjamini & Hochberg, 1995). We also used bootstrapping to get 95% confidence intervals (CI) of each mean difference. Lastly, we supplement our pairwise tests with Bayes factors (BFs), which quantify the strength of evidence (Rouder et al., 2009). The BF is the ratio of the probability of the data under the alternate hypothesis (that two means differ) relative to the probability of the data under the null hypothesis (that there is no difference). These ratios are sometimes called BF₁₀: a BF of 10 would indicate that the data are 10 times more likely under the alternate hypothesis than the null. BFs between 3 and 10 are regarded as substantial evidence for the alternate hypothesis, and BFs greater than 10 as strong evidence. Conversely, BFs between 1/3 and 1/10 are considered substantial evidence for the null hypothesis, etc. We computed BFs using the bayesFactor MATLAB toolbox (Krekelberg, 2024)

We provide both p-values and BFs so that the reader may be fully informed to judge the strength of evidence. We make strong claims about rejecting a hypothesis only when a test that yields both a low p-value and a high BF.

To compare across tasks across experiments, we also fit linear mixed effect (LME) models to single-trial data, with fixed effects of the task, the set size (number of words), random intercepts and slopes by participant, and random effects for individual stimulus items. We use these LMEs to compare redundant target effects across tasks, and they yield F statistics similar to a two-way repeated measures ANOVA (but with degrees of freedom that take into account all the individual trials). The p-values for these tests were corrected for false discovery rate across all four experiments.

Estimates of the redundant target effect are complicated by the possible differences in performance between the two single-target conditions (Mullin et al., 1988). Such differences are the rule in multisensory (e.g. auditory-visual) experiments, but even occur with two visual stimuli at isoeccentric locations. For example, in the current experiments, single words were judged more accurately and quickly at the top location than at the bottom location (see results below).

Without such a difference between locations, the redundant target effect can be simply estimated with the “averaging baseline” approach: the difference between the two-target condition and the average of the one-target conditions (averaging across locations). But with a known difference between locations, the averaging approach can overestimate the size of the redundant target effect. An alternative is the “fastest baseline” approach: to estimate the effect as the difference between the two-target trials and trials with one target at the location that produces *faster* responses in the mean. This approach reduces the estimated magnitude of a positive redundant target effect. It is therefore more conservative for the purpose of estimating the *positive* redundant target effects predicted by many parallel models. Some authors recommend an alternate approach to avoid positive biases in the redundant target effect estimation: the baseline

is set for each subject to their individual fastest location (if they have a significant bias) (Miller & Lopes, 1988).

The goal of the current study is to investigate possible *negative* redundant target effects (slower responses to 2 targets than to 1 target). If we were to compare one-target RTs at the fastest location to two-target RTs, we would *overestimate* the magnitude of a negative effect. Thus, we use the averaging baseline throughout this article because it is more conservative for detecting negative effects.

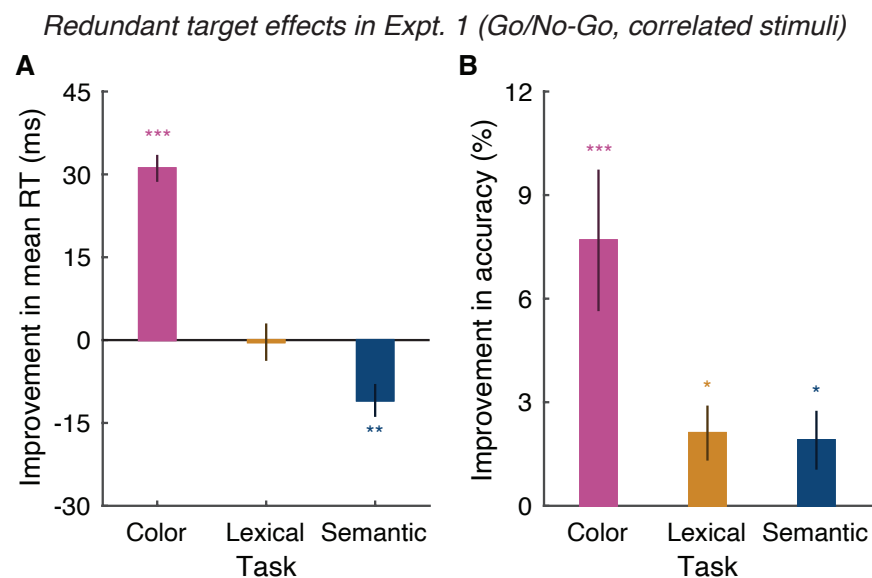


Figure 4: Redundant target effects in Experiment 1. These bar plots show the mean *improvement* in (A) mean correct response time (RT) and in (B) accuracy for 2 targets compared to 1 target. The mean performance levels from which these difference scores were derived are in Figure 5. Error bars are ± 1 SEM. Asterisks indicate that the mean effect is significantly different from 0 (** $p < 0.001$, ** $p < 0.01$, * $p < 0.05$, FDR-corrected).

Results

Response times: **Figure 4A** shows that in the color task, there was a positive redundant target effect: a speeding of correct responses to two targets compared to one, by 31 ms on average. However, in the semantic task, there was a significantly negative effect (a slowing of responses) of 11 ms on average. In the lexical decision task, there was

no effect of redundant targets. **Table 3** lists the statistics on each effect. The mean response times in each individual condition (rather than the differences between one and two targets) are shown in the top row of **Figure 5**.

To compare the redundant target effects across tasks, we also fit three linear mixed-effect models to single-trial correct response times (target-present trials only). Compared to the color task, both the lexical and semantic tasks had significantly different redundant target effects (both $F > 36$, $p < 10^{-8}$, $BF > 10^6$). The lexical and semantic tasks yielded effects that were not significantly different ($F(1, 15475) = 3.85$, $p = 0.050$, $BF = 2.3$).

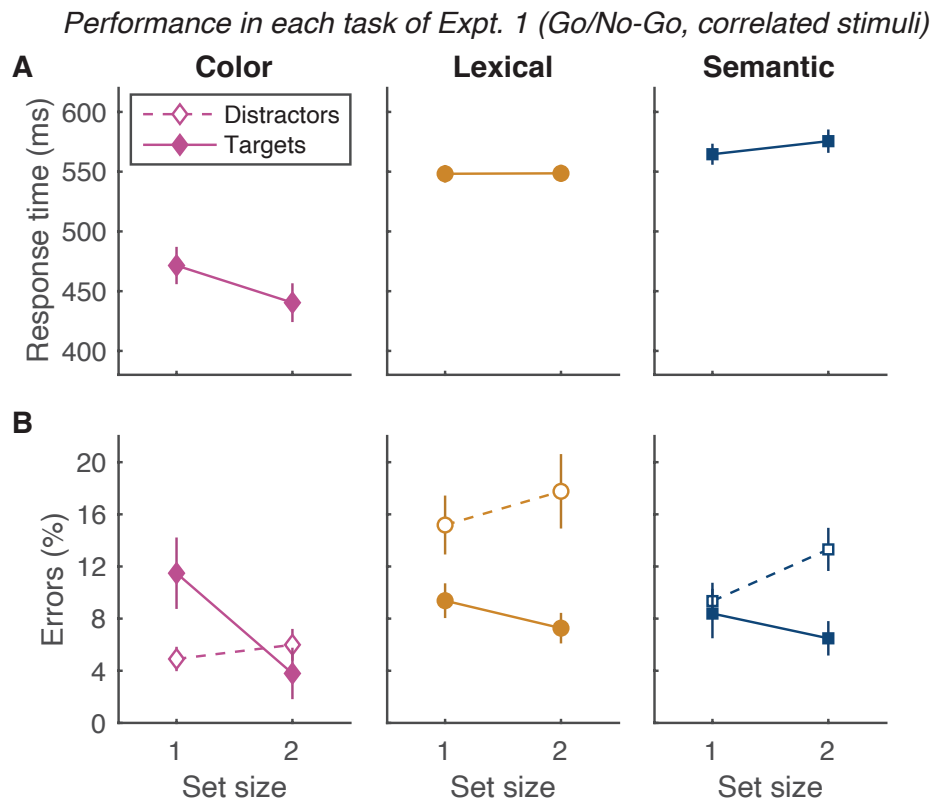


Figure 5: Mean performance in each task of Experiment 1, plotted separately for targets and distractors, for set size 1 vs. 2. **(A)** Mean correct response times. Note that there are no correct response time data for distractors, because in these go/no-go tasks, the correct response to distractors was to not press any key. **(B)** Percent of trials with errors. Data for distractors are plotted with open symbols and dashed lines (showing how often the participants made false alarms). Error bars are ± 1 SEM.

Accuracy: On average, collapsing across all conditions, participants achieved 93%, 88%, and 91% correct accuracy in the color, lexical, and semantic tasks, respectively (SEMs = 1, 2, and 1%). **Figure 4B** shows the mean improvements in accuracy caused by a redundant target compared to trials with a single target. There was a significant improvement in all three tasks (fewer misses; see statistics in Table 3). The mean percentage of trials with errors in each condition (including for trials with distractors) are plotted in the bottom row of **Figure 5**. Compared to the color task, both the lexical and semantic tasks had smaller redundant target effects on accuracy for detecting targets (both $F > 24$, $p < 10^{-6}$, $BF > 3$). The effects did not differ significantly between the lexical and semantic tasks ($F(1, 16776) = 0.75$, $p = 0.39$, $BF = 0.86$).

Task	1 Targ. Mean	2 Targ. Mean	Effect Mean	Effect SEM	95% CI	<i>t</i>	<i>p</i>	BF
<i>Correct response time (ms)</i>								
Color	471.5	440.4	31.1	2.4	[26 36]	12.50	6.75×10^{-12}	1.37×10^{10}
Lexical	548.2	548.6	-0.4	3.4	[-7 7]	-0.11	0.915	0.20
Semantic	564.6	575.5	-10.9	3.0	[-17 -5]	-3.60	0.002	27.48
<i>Errors (percent)</i>								
Color	11.48	3.79	7.69	2.05	[4.49 12.61]	3.69	0.001	34.35
Lexical	9.38	7.27	2.11	0.80	[0.75 3.88]	2.60	0.018	3.26
Semantic	8.39	6.49	1.90	0.85	[0.51 4.03]	2.19	0.038	1.54

Table 3: Statistics on redundant target effects in Experiment 1, for correct response times (top three rows) and error rates (bottom three rows). The columns “1 Targ. Mean” and “2 Targ. Mean” refer to the mean response times and error rates on 1- and 2-target trials, respectively. The other columns list statistics on the redundant target effect, calculated as the mean *improvement* in response times and error rate, contrasting 1-target displays (averaged over sides) vs. 2-target displays. Note that the errors here are all misses (misclassifying a target as a distractor). The degrees of freedom for the t-tests was 27. For each measure (response time or accuracy), p-values are corrected for false discovery rate across all 12 comparisons including all 4 experiments in the study. BF = Bayes Factor.

Differences between top and bottom sides: on single-target trials, participants generally made faster correct responses and fewer errors for targets at the top location compared to the bottom location. Correct RTs were on average 30, 33, and 15 ms faster

for the top location in the color, lexical, and semantic tasks, respectively (SEMs = 5 ms, all FDR-corrected $p < 0.01$, $BFs = 191$, 2×10^4 , and 9, respectively). Error rates were on average 5, 8, and 2% lower at the top location in those three tasks, respectively (SEMs = 4, 2, and 1%; only significant in the lexical task, $p < 0.0001$, $BF = 2514$).

As discussed in the Analysis section above, an alternate calculation of the redundant target effect is the difference between trials with two targets and trials with one target at the top location only, because RTs tend to be faster for words at the top. This shifts estimates of the mean redundant target effect down, rendering even the lexical task's effect significantly negative (mean = -16 ms, $p = 0.001$) while maintaining a significantly positive effect in the color task (mean = 23 ms, $p = 1 \times 10^{-8}$). Given that we are most interested in detecting the negative effects predicted by our generalized serial model, in the primary analyses above (Table 3) we use the approach that is more conservative for negative effects: comparing two-target responses to the *average* of responses to single targets at the top and bottom locations.

Discussion

The redundant target effects in the first experiment suggest that the colors of the letters within two words can be processed in parallel, leading to speeding of response times. However, the meanings of the two words are not necessarily processed in parallel. This is because the presence of a second word target in the lexical decision task yielded 0 improvement in response time, and the semantic categorization task yielded a significant slowing of response time.

Our generalized models show that such a negative response time effect is consistent with the standard serial model, even when accompanied by an increase in accuracy (see Appendix). It can be explained by participants occasionally miscategorizing the first target they process as a distractor, then going on to correctly process the other target, with a slower response time compared to correct trials with single

targets. The negative effect in the semantic task is also consistent with some (but not all) fixed-capacity parallel models. One potentially interesting result is that the negative redundant target effect was significant in the semantic task but not the lexical task. However, we lack strong statistical evidence that those two results were significantly different from each other.

In theory, the comparison of two-target trials to one-target trials could be influenced by sensory interactions between the two simultaneous targets, such as crowding. However, the words were presented on opposite sides of fixation at distances too great for crowding (Pelli & Tillman, 2008). Another possible concern is that participants simply are not able to divide their attention between those two locations simultaneously. However, the positive redundant target effect in the color task rules out that concern as an explanation for the non-positive effects in the lexical and semantic tasks.

It is also noteworthy that redundant targets improved accuracy in all three tasks (although that effect was significantly larger in the color task than in the other two). Is that evidence of parallel processing of two words in all three tasks? Not necessarily: the serial model also predicts improvements in accuracy that go along with slowing of response speeds. This is merely a statistical facilitation: there are two chances to reach the correct decision when two targets are present. Also note that, as shown in the bottom row of Figure 5, errors on trials with distractors *increase* when the set size is 2 compared to 1 (a decrease of accuracy, opposite to the pattern for trials with targets). This might also be consistent with a shift in decision criterion, as participants are somewhat more likely to report “target present” when they see two words than when they see one word.

The color task was overall easier than the other two tasks, with faster RTs and lower error rates. Experiment 4 addresses this issue by making the color task more difficult than the other two tasks. As shown below, the conclusions did not change.

Thus, the entire data set allows us to rule out the standard serial model only for the color task, which yielded a positive redundant target effect in response times as well as accuracy. The semantic task shows a negative effect on response times, which is quite rare in the redundant target literature and can be readily accounted for by the standard serial model.

Experiment 2: Choice procedure with correlated stimuli

Experiment 1 assessed redundant target effects with the simplest possible design (“correlated stimuli”, meaning no trials with mixed targets and distractors) and the procedure thought to be most sensitive (go/no-go). In the next two experiments, we used variations of the paradigm that have previously been used to study word recognition and might alter the participant’s strategy. In Experiment 2, we used the same “correlated” stimulus conditions as Experiment 1, but we required participants to make a choice response (target present vs absent) on every trial.

Method

Participants: Participants were recruited in the same way as in Experiment 1. See Table 1 for counts. Applying the same accuracy criteria as in Experiment 1, no participants had to be excluded.

Stimuli and procedure: All methodological details were the same as in Experiment 1, except the participant had to make a categorization judgment on every trial: press the left arrow if no target was present on the screen, or the right arrow if any targets were present on the screen. The response interval was unlimited, but participants were requested to “respond as quickly as you can without making unnecessary errors.”

Results

Response times: **Figure 6A** demonstrates that the redundant target effects on correct response times in Experiment 2 were similar to those in Experiment 1. There was a significant speeding of response times in the color task (by 27.7 ms on average), no significant effect in the lexical task (mean = -5.1 ms), and a significant *slowing* in the semantic task (by -17.5 ms). Statistics on the effect for each task are reported in **Table 4**. Compared to the color task, both the lexical and semantic tasks had significantly different redundant target effects (both $F > 18$, $p < 10^{-4}$, $BF > 10^3$). The lexical and semantic tasks yielded effects that were not significantly different ($F(1, 15141) = 2.61$, $p = 0.11$, $BF = 0.97$). See the top row of **Figure 7** for mean correct response times in each condition separately.

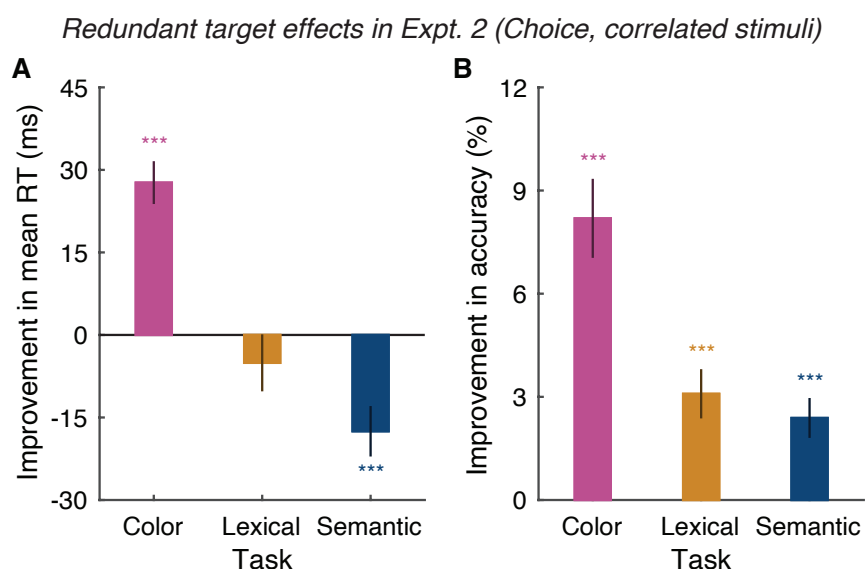


Figure 6: Mean redundant target effects in Experiment 2. Format as in Figure 4.

Accuracy: On average, collapsing across all conditions, participants achieved 94%, 90%, and 92% correct accuracy in the color, lexical, and semantic tasks, respectively (SEMs = 1%). **Figure 6B** plots the mean improvements in accuracy (fewer misses) caused by redundant targets, which were significant in all three tasks (as also reported in Table 4). The mean percent errors in each condition are plotted in the bottom row of **Figure 7**. Compared to the color task, both the lexical and semantic tasks had smaller redundant

target effects on accuracy (both $F > 21$, $p < 10^{-5}$, $BF > 50$). The redundant target effects did not differ significantly between the lexical and semantic tasks ($F(1, 16713) = 0.03$, $p = 0.85$, $BF = 0.34$).

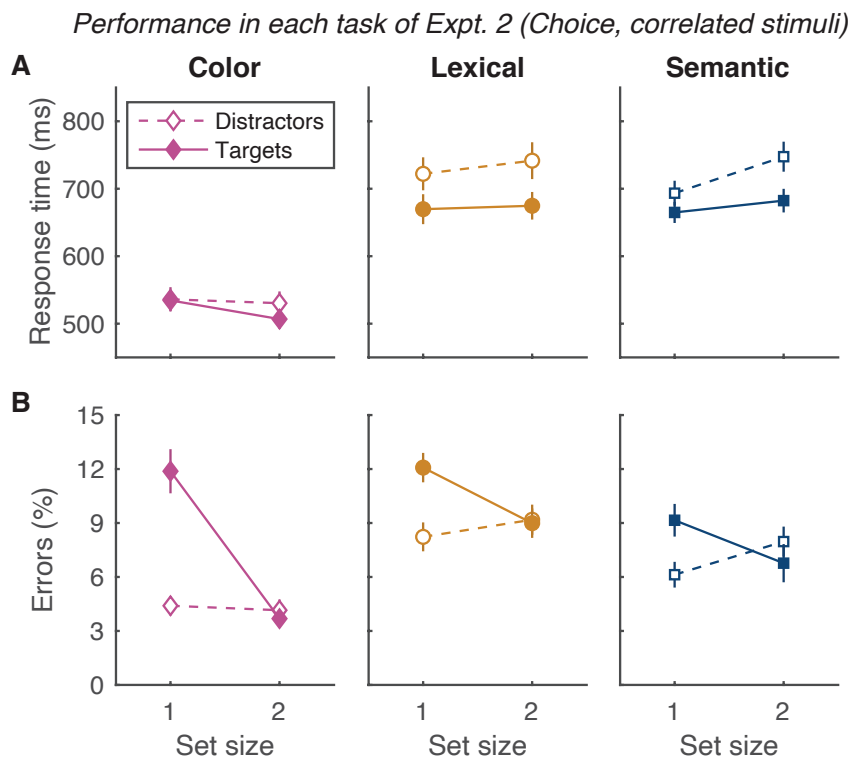


Figure 7. Mean performance in each task of Experiment 2. Format as in Figure 4.

Task	1 Targ. Mean	2 Targ. Mean	Effect Mean	Effect SEM	95% CI	<i>t</i>	<i>p</i>	BF
<i>Correct response time (ms)</i>								
Color	534.5	506.8	27.7	3.9	[20 35]	6.99	5.3x10 ⁻⁷	115,304
Lexical	669.7	674.8	-5.1	5.2	[-15 5]	-0.96	0.3757	0.31
Semantic	664.9	682.4	-17.5	4.6	[-29 -10]	-3.75	0.0011	38.7
<i>Errors (percent)</i>								
Color	11.88	3.69	8.19	1.15	[6.18 10.88]	7.00	1.4x10 ⁻⁶	118,545
Lexical	12.08	8.99	3.09	0.72	[1.60 4.49]	4.24	0.0006	121.4
Semantic	9.15	6.77	2.39	0.58	[1.26 3.52]	4.04	0.0008	76.0

Table 4: Statistics on redundant target effects in Experiment 2, formatted as in Table 3. The degrees of freedom was 28 for the color task and 27 for the others.

Differences between top and bottom sides: Correct RTs were on average 35, 36, and 31 ms faster for single targets the top location in the color, lexical, and semantic tasks, respectively (SEMs = 7, 10 and 6 ms, all $p < 0.01$, all $BF > 35$). Error rates were 8, 7, and 3% lower for single targets at the top location in those three tasks, respectively (SEMs = 2, 1, and 1%, all $p < 0.01$, all $BF > 13$).

Discussion

The results of Experiment 2, which used a choice procedure, were consistent with the results of Experiment 1, which used a go/no-go procedure. The redundant target effects on response times were consistent with parallel processing in the color task and serial or fixed-capacity parallel processing in the lexical and semantic tasks. Again, the redundant target effect in the semantic task was significantly negative.

Experiment 3: Forced-choice procedure with uncorrelated stimuli

In both experiments reported so far, the words presented simultaneously on trials with set size 2 were always of the same category (both targets or both distractors, although never the same exact words). In other words, there was a contingency between the category of the stimulus at one location and the category of the other. That is what we mean by a ‘correlated’ stimulus design. One potential drawback of this design is that the participant might adopt a strategy of always processing just one word, knowing that the other word leads to the same correct decision. (Although they cannot simply pick a *side* of the screen and always ignore stimuli presented on the other side, because half the trials contain just a single word that could be on either side, unpredictably). In the third experiment, we addressed this issue by including trials in which a target is paired with a distractor. Thus, the stimuli are uncorrelated. This should motivate the participant to process both stimuli as well as they can. Uncorrelated stimuli like this are also common in the wider visual search literature.

The “uncorrelated” design used in Experiment 3 (and 4) also differs from Experiments 1-2 in the “interstimulus contingencies” (Mordkoff & Yantis, 1991). These contingencies describe the probability of a target at one location contingent on what is presented at the other location. Note that within each experiment the contingencies are consistent across all three tasks. We take this issue up in the General Discussion.

Method

Participants: Participants in the lexical and semantic tasks were recruited and compensated in the same way as in Experiment 1, via Prolific. The target sample size was maintained at 28. Participants in the color task were recruited from the Barnard College Introductory Psychology subject pool, and participated in exchange for course credit. See Table 1 for counts of included and excluded participants.

		Bottom word		
		<i>None</i>	<i>Distractor</i>	<i>Target</i>
Top word	<i>None</i>	<i>N/A</i>	0.05	0.15
	<i>Distractor</i>	0.05	0.15	0.15
	<i>Target</i>	0.15	0.15	0.15

Table 5: The probability of stimulus pairings at the top and bottom locations in Experiment 3. The word at each location was either absent, a distractor, or a target. The green shading highlights conditions when 2 words were present. In this design, the two stimuli were uncorrelated and independent: the conditional probability of one stimulus being a target given that the other was a target was 0.5.

Stimuli and procedure: All details were the same as in Experiment 2, except as noted here. The primary difference is that 30% of trials contained mixed pairs of 1 target and 1 distractor. **Table 5** lists the proportions of trials assigned to each type, which were

chosen to ensure that the categories (target vs distractor) of the upper and lower stimuli were independent (uncorrelated). Specifically, on two-word trials, the conditional probability of one stimulus being a target given that the other was a target was 0.5. In contrast, this conditional probability was 1.0 in Experiment 1 and 2. Another difference in Experiment 3 was that across the experiment, the probability of a target being present on any given trial was 0.75, rather than 0.5. That was true both among trials with set size 1 and trials with set size 2. To maintain the same number of two-target trials as in Experiments 1 and 2, we increased the total number of trials in Experiment 3 to 1000. Each participant conducted 10 blocks of 100 trials each. Also, six words were added to the stimulus set for the color and semantic tasks (see the online data repository).

Redundant target effects in Expt. 3 (Choice task, uncorrelated stimuli)

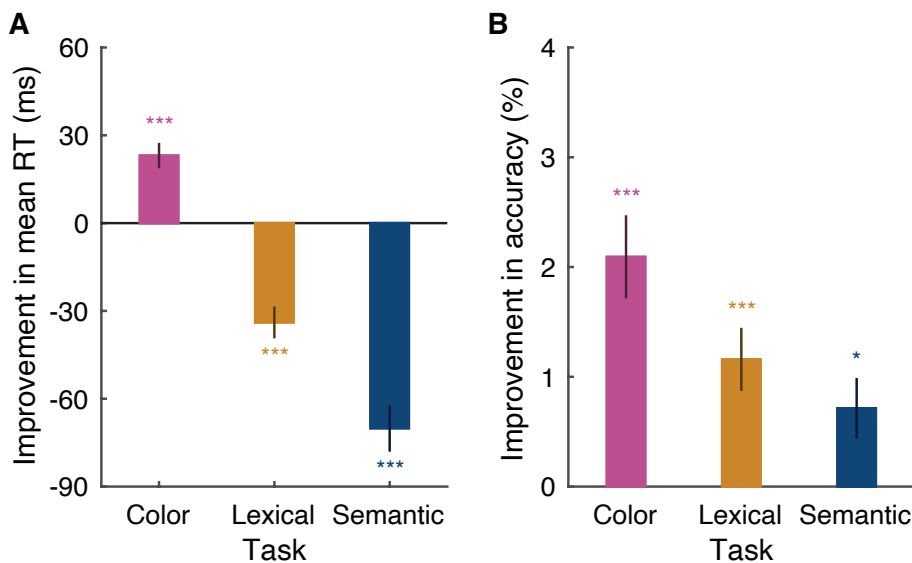


Figure 8: Redundant target effects in Experiment 3. Format as in Figures 4 and 6. These bar plots show the mean *improvement* in correct response time (RT) and accuracy comparing trials with 2 targets to trials with just 1 target and 0 distractors. The mixed-pair trials were not included in this analysis (but are plotted in Figure 9).

Results

Response times: **Figure 8A** shows that the redundant target effects were consistent with the prior two experiments, although somewhat magnified. As detailed in **Table 6**,

redundant targets improved responses in the color task (by 23 ms on average), but slowed responses in *both* the lexical task (-34 ms) and the semantic task (-70 ms). All three pairwise comparisons between these effects were significant (color vs. lexical and color vs. semantic: both $F > 27$, $p < 10^{-6}$, $BF > 10^6$; lexical vs semantic: $F(1, 24587) = 5.70$, $p = 0.022$, $BF = 42$). See the top row of **Figure 9** for mean correct response times in each condition separately, including the mixed target-distractor trials, which are represented with single lightly-shaded symbols.

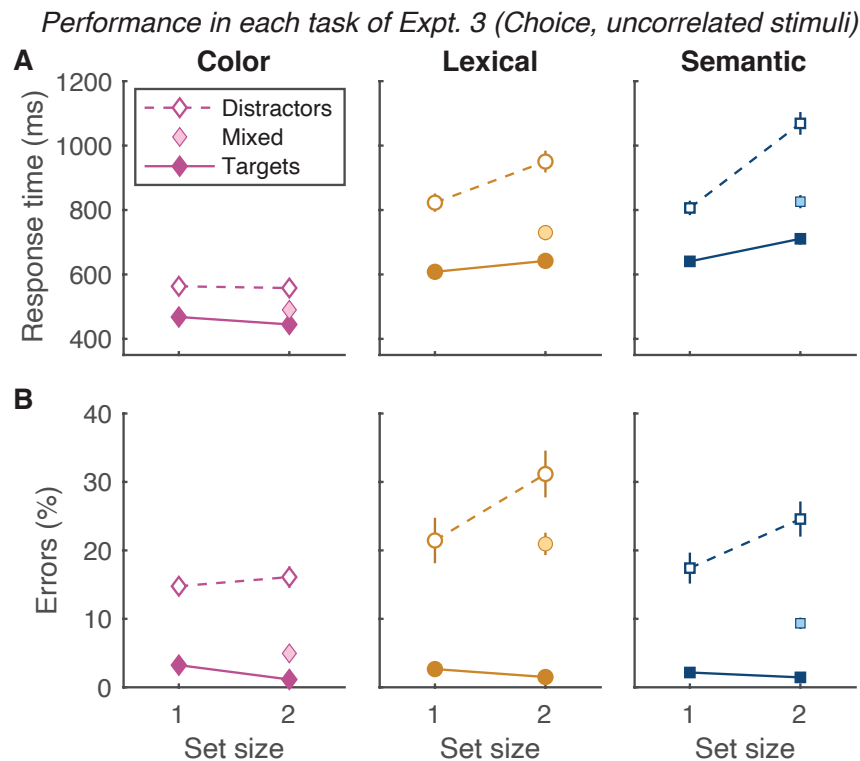


Figure 9. Mean performance in each task of Experiment 3. Format as in Figures 5 and 7, except that this experiment included ‘mixed’ trials in which 1 target appeared with 1 distractor. That mixed-pair condition is represented by the single symbol with medium fill in each panel.

Accuracy: On average, collapsing across all conditions, participants achieved 93%, 86%, and 92% correct in the color, lexical, and semantic tasks, respectively (SEMs = 1%).

Figure 8B plots the positive effects of redundant targets on accuracy (fewer misses), which were significant in all three tasks (see Table 6 for a summary of target-present

trials). The effect on accuracy was significantly larger in the color task than the semantic task ($F(1,24907)=4.90$, $p=0.046$, $BF=7.5$), but the other two pairwise comparisons across tasks were not significant (color vs. lexical: $F(1, 24041)=2.71$, $p=0.15$, $BF=1.2$; lexical vs semantic: $F(1, 24145)=0.65$, $p=0.46$, $BF=0.45$). The mean percentage errors in each condition are shown in the bottom row of **Figure 9**. Note that there is a larger difference between accuracy for targets (hit rates plotted as solid lines in Fig. 9B) and accuracy for distractors (false alarm rates plotted as dashed lines) in this experiment than in the prior experiments. This can be explained by the fact that targets were present on 75% of trials overall, so participants were more liberal in reporting “yes” and made more false alarms on distractor trials. This was a consequence of making the two stimulus categories independent (see Table 5 above).

Task	1 Targ. Mean	2 Targ. Mean	Effect Mean	Effect SEM	95% CI	<i>t</i>	<i>p</i>	BF
<i>Correct response time (ms)</i>								
Color	467.9	444.8	23.1	4.3	[15 32]	5.25	3.2×10^{-5}	1,409
Lexical	608.2	642.1	-33.9	5.4	[-44 -24]	-6.11	5.3×10^{-6}	8,640
Semantic	640.6	710.7	-70.1	8.0	[-88 -56]	-8.59	2.0×10^{-8}	3.6×10^5
<i>Errors (percent)</i>								
Color	3.24	1.14	2.09	0.38	[1.33 2.82]	5.42	3.96×10^{-5}	2,150
Lexical	2.66	1.50	1.16	0.29	[0.56 1.72]	3.95	0.0010	56.7
Semantic	2.16	1.45	0.71	0.28	[0.19 1.26]	2.54	0.0189	2.9

Table 6: Statistics on redundant target effects in Experiment 3, formatted as in Table 3. This table only reports performance on trials with 2 targets as compared to trials with 1 target (trials with distractors are not included here). The degrees of freedom for the t-tests was 27 for the color and semantic tasks, and 25 for the lexical task.

Differences between top and bottom sides: Correct RTs were on average 20, 32, and 25 ms faster for single targets at the top than bottom location in the color, lexical, and semantic tasks, respectively (SEMs = 5, 9 and 6 ms, all $p < 0.01$, $BF > 14$). Error rates were on average 2.8, 0.3, and 0.9% lower for single targets at the top location in those three tasks,

respectively (SEMs = 0.5, 0.4, and 0.3%, significant in the color and semantic tasks, $ps < 0.02$, $BFs > 6$).

Discussion

The results of Experiment 3 were consistent with both prior experiments: there is a positive redundant target effect on response times only in the color task. One new result in this experiment was that the lexical task, as well as the semantic task, yielded a significantly *negative* redundant target effect (slowing of responses). These results again reject the standard serial model for the color task. The serial model can account for lexical and semantic task performance, as can some versions of fixed-capacity parallel models.

In order to test the prediction of the standard serial model, focused on comparing two-target trials (set size 2) to single-target trials (set size 1). But as shown in Figure 9, for all three tasks, mean responses on two-target trials were faster (and more accurate) than responses on mixed-pair trials (with 1 target and 1 distractor). As explained in the Introduction, the standard serial model predicts this effect: on some mixed-pair trials, the non-target is processed first and then search continues to correctly respond to the target, with a slow response time. The parallel models also predict this effect. Thus, the mixed-pair trials are not useful for distinguishing these models.

Experiment 4: Forced-choice procedure with uncorrelated stimuli on the left and right

In Experiments 1-3, the words were positioned just above and below fixation. We made this choice to minimize the distances of all the letters from fixation and to match the design of Mullin & Egeth (1989). But English words are typically arranged horizontally. It is possible that parallel processing is strongest when words are presented in that standard arrangement.

Accordingly, Experiment 4 used the same tasks and experimental design as Experiment 3, except the words were arranged horizontally, immediately to the left and

right of the point of fixation, with a single blank letter space between their inner letters. We also increased the difficulty of the color task by reducing the saturation of the colored letters. This addresses a concern that in the prior experiments the color task was easier than the other two tasks.

Method

Participants: Participants were recruited and compensated in the same way as in Experiment 1, via Prolific. See **Table 1** for the counts of included and excluded participants.

Stimuli and procedure: All details were the same as in Experiment 3, except as noted here. An example stimulus display from the color task is shown in **Figure 10A**. The words were placed to the left and/or right of fixation, centered on the horizontal midline. The empty space between the words was always one letter space in width. There was no central fixation cross. Instead, the point of gaze fixation was marked with two thin black vertical lines centered on the screen's vertical midline. The length of each bar was 8% of the total screen height. The top bar's lower end was roughly in line with the top of the tallest letter, and the bottom bar's upper end was the same distance from the horizontal midline. Participants were instructed to fixate on the space between the two bars.

We also made the color task more difficult than in the previous experiments by reducing the relative saturation of the color targets by more than half. On a scale from 0 to 255, the RGB values of the red letters were (92, 59, 59), which corresponds to 36% saturation, relative to the maximum possible on the screen. The RGB values of the green letters were (61, 88, 61), which corresponds to 31% saturation. See an example colored word in Figure 10A.

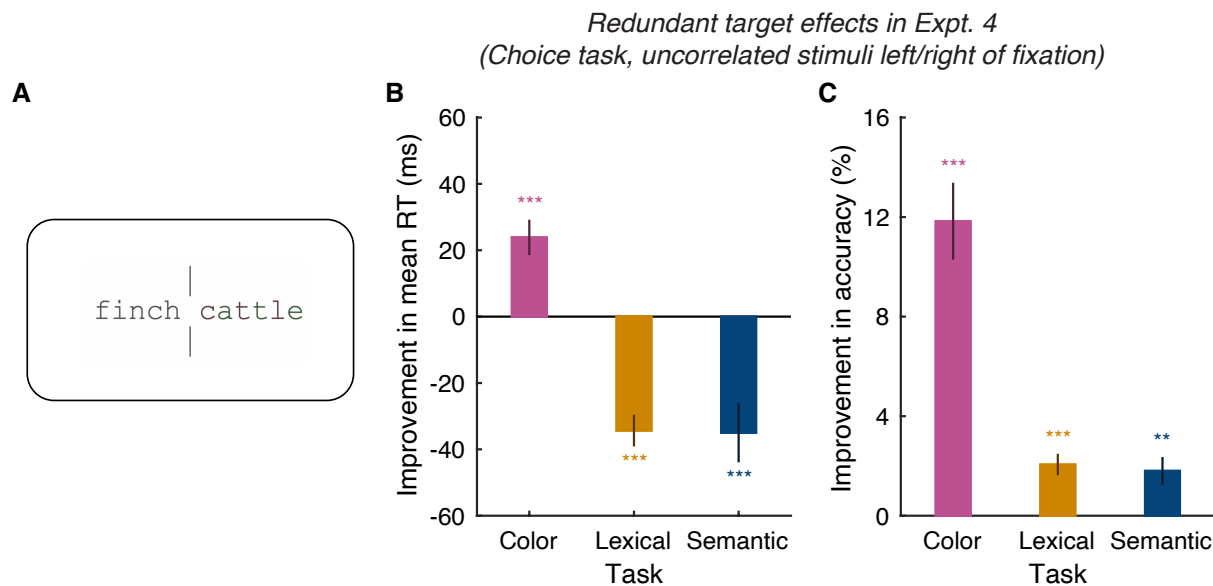


Figure 10: (A) Example stimuli in Experiment 4, showing a gray word on the left and a colored word on the right (with lower saturation than in Experiments 1-3). (B) and (C) Redundant target effects in Experiment 4. Format as in Figure 8.

Results

Response times: **Figure 10B** shows that the redundant target effects were consistent with the other three experiments, despite the change in stimulus positions. As detailed in **Table 7**, redundant targets improved responses in the color task (by 24 ms on average), but slowed responses in both the lexical task (-34 ms) and the semantic task (-35 ms). The redundant target effect was significantly different in the color task as compared to the lexical task and to the semantic task (both $F > 26$, $p < 10^{-6}$, $BF > 4000$). The effects in the lexical and semantic tasks did not differ ($F < 0.001$, $p = 0.98$, $BF = 0.28$). See the top row of **Figure 11** for mean correct response times in each condition separately.

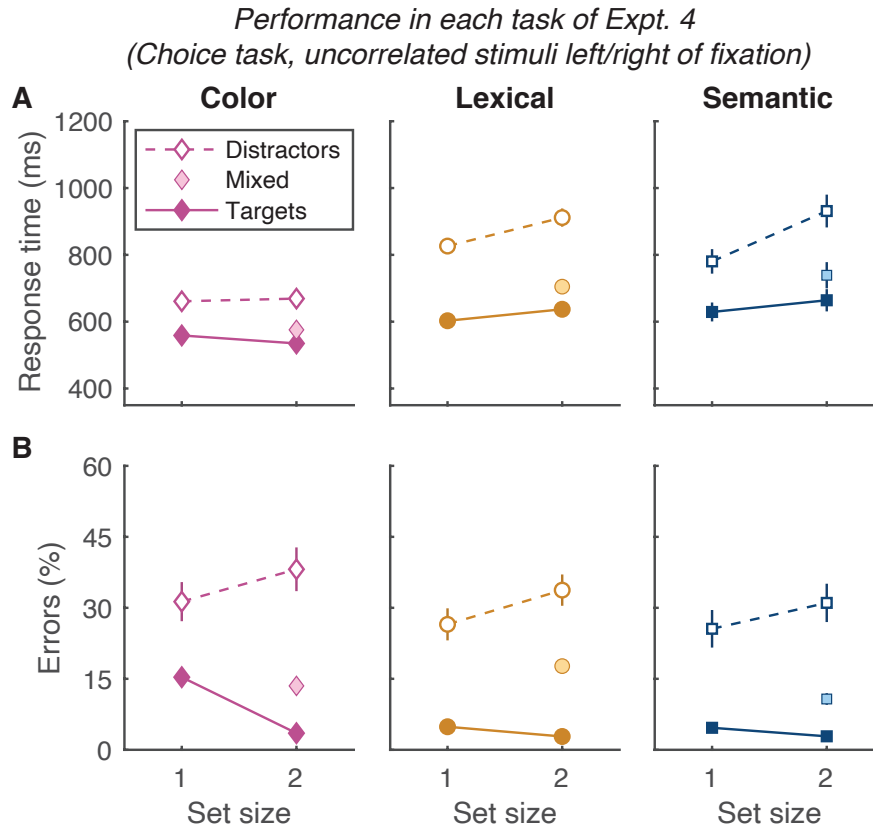


Figure 11. Mean performance in each task of Experiment 4. Format as in Figure 9.

Accuracy: On average, collapsing across all conditions, participants achieved 82%, 84%, and 88% correct in the color, lexical, and semantic tasks, respectively (SEMs = 2%, 1% and 1%). Thus, unlike in the previous three experiments, the color task was now the hardest of the three tasks, likely because the saturation of the colored letters was much lower. **Figure 10C** plots the positive effects of redundant targets on accuracy (fewer misses), which were significant in all three tasks (see Table 7 for a summary of target-present trials). The effect on accuracy was significantly bigger in the color task than in the other two tasks (both $F > 12$, $p < 0.001$, $BF > 10^5$), and did not differ between the lexical and semantic tasks ($F = 0.67$, $p = 0.46$, $BF = 0.29$). The mean percentage errors in each condition are shown in the bottom row of **Figure 11**. As in Experiment 3, which also used an uncorrelated stimulus design with targets present on 75% of all trials, error rates were higher on trials with distractors than with targets.

Task	1 Targ. Mean	2 Targ. Mean	Effect Mean	Effect SEM	95% CI	<i>t</i>	<i>p</i>	BF
<i>Correct response time (ms)</i>								
Color	558.7	534.9	23.9	5.3	[12 33]	4.38	0.0004	117
Lexical	602.8	637.1	-34.3	4.8	[-43 -25]	-7.05	6.5×10 ⁻⁷	73,980
Semantic	629.1	664.1	-35.0	8.9	[-55 -19]	-3.85	0.0011	45.6
<i>Errors (percent)</i>								
Color	15.34	3.51	11.83	1.55	[9.00 14.89]	7.48	1.4 ×10 ⁻⁶	66,861
Lexical	4.87	2.81	2.06	0.43	[1.23 2.90]	4.68	0.0003	308
Semantic	4.65	2.85	1.80	0.56	[0.94 3.34]	3.17	0.0054	10.2

Table 7: Statistics on redundant target effects in Experiment 4, formatted as in Table 6. This table only reports performance on trials with 2 targets as compared to trials with 1 target (trials with distractors are not included here). The degrees of freedom for the t-tests was 21 for the color task and 25 for the lexical and semantic tasks.

Differences between the left and right sides: Correct RTs were on average 38, 64, and 47 ms faster for single targets at the right than left location in the color, lexical, and semantic tasks, respectively (SEMs = 9, 8 and 8 ms, all $p < 0.001$, $BF > 65$). Error rates were on average 5, 3, and 2% lower for single targets at the right location in those three tasks, respectively (SEMs = 3, 1, and 1%). That effect was not significant in the color task ($p = 0.21$, $BF = 0.57$), but it was in the lexical and semantic tasks, $ps < 0.01$, $BFs > 10$). A right visual field advantage for word recognition has been observed many times (reviewed by Yeatman & White, 2021).

Discussion

The results of Experiment 4, with horizontally arranged words, were qualitatively consistent with the results of Experiment 3. Both the lexical and semantic tasks produced negative redundant target effects, while the color task produced a positive redundant target effect despite now being the most difficult of the three tasks. We conclude that the

negative redundant target effects in Experiments 1-3 cannot be accounted for by the unnatural vertical arrangement of the words above and below fixation.

General Discussion

The four experiments reported above consistently demonstrate that there is a positive redundant target effect when the targets are defined by color but not when the targets are defined by lexicality or by semantic category. In all three tasks, the stimuli were written words presented singly or in pairs near fixation. In the color task, the presence of a redundant target (a word written in colored letters) consistently sped responses compared to trials with a single target. In the semantic task, the presence of a redundant target (a word that refers to a living thing) significantly *slowed* responses in all of the experiments. In the lexical decision task, the redundant target had no significant effect in two experiments (with correlated stimuli), and a significantly negative effect in the third and fourth experiments (with uncorrelated stimuli).

In this study, the color task served as a positive control condition to demonstrate the expected redundancy gain for a simple visual feature task. Indeed, the prior literature is dominated by reports of positive redundant target effects (Townsend, 1990; van der Heijden et al., 1983). The positive effects we found with the color task in all four experiments demonstrate that participants can attend to words at the chosen positions. What we highlight is the existence of a negative redundant target effect with nearly identical stimuli at the same positions, but when the task requires semantic categorization of the words' meanings. Such a negative effect is quite rare and can be explained by our updated models that incorporate errors.

In Experiments 1-3, the color task was overall easier (lower error rates) than the lexical and semantic tasks. But in Experiment 4, we reduced the salience of the color targets to make the color task the most difficult of the three. The positive redundant target

effect persisted. Thus, any differences in overall difficulty cannot explain the different results across tasks.

These data are all consistent with the hypothesis that the low-level features of multiple stimuli, such as their color, are processed in parallel with little cost, but there is a capacity limit for recognizing the meanings of multiple written words at once. The standard serial model predicts the negative redundant target effect for semantic judgements, as long as there are some errors in classifying targets, and thus cannot be rejected. Alternatively, some fixed-capacity parallel models can also account for the negative effect.

Relation to previous redundant target studies of word recognition

In contrast to our key result, some prior studies have reported positive redundant target effects for word recognition. However, most of those experiments presented two copies of the same word on redundant target trials (Egeth et al., 1989; Hasbrooke & Chiarello, 1998; Mohr et al., 1994, 1996; Mullin & Egeth, 1989). The resulting redundant target effects might be explained by facilitation or coactivation at the stage of letter or syllable processing (as discussed in the “Interactive Processing” section below). The only study to report a positive redundant target effect that did not present identical pairs of words was by Shepherdson & Miller (2014). They found that semantic categorization judgments were faster for two targets than for a single target paired with a pronounceable non-word. This could be interpreted as a positive redundant target effect and evidence for parallel processing. However, the result can be explained by the standard serial model if we assume that on some one-word trials, the participant processes the non-word before they process the target. Therefore, we consider the strict test of our serial model to be the contrast between trials with two targets (which are two different words) and trials with a single target presented alone.

Thus far, all experiments that made this strict test of the serial model with word recognition tasks lead to the same result. They were conducted by us in the present study and by Mullin & Egeth (1989). The experiments in that prior study were like our Experiment 1: they presented words above and/or below fixation and used go/no-go target detection tasks in which targets were never presented with distractors (a correlated stimulus design). In two of their experiments, the words presented together on redundant-target trials were *not* identical (as in ours). In those experiments, they found that both lexical decision and semantic categorization judgments were *slowed* by the presence of a redundant target – but significantly so only in the lexical task.

Based on these results, the authors rejected the standard, self-terminating, unlimited-capacity parallel processing model for recognizing two words, as do we. They raised several tentative explanations for the negative redundant target effect in the lexical task. On the one hand, they hypothesized some type of ‘interference’ between two words that are processed in parallel. This is similar to the newer ideas about orthographic interference that are discussed below. On the other hand, they argued that serial processing could explain the effects if both words were processed exhaustively, that is, if search was not self-terminating. But our model generalizations presented above make clear that the negative redundant effect is predicted by the standard self-terminating serial model as long as for errors in stimulus classification are taken into account.

We also go beyond previous studies by showing that the results are consistent whether the task requires a go/no-go response or a choice response, and whether targets can appear with distractors, and whether the words are presented above and below fixation or to the left and right. The negative effect was always negative for the semantic task. For the lexical task it was never positive, and significantly negative in the two experiments that used the uncorrelated stimulus design.

Thus, considering both our data and Mullin & Egeth (1989), we can reject the standard self-terminating serial model for the color task, but not for the word recognition tasks (i.e., lexical decision and semantic categorization). In addition, we can reject the standard, self-terminating, unlimited-capacity parallel model for the word recognition tasks, but not for the color tasks. What we cannot do is reject the standard, self-terminating *fixed-capacity* parallel model for the word recognition tasks. That does not mean that the fixed-capacity model is always viable. It can predict the wide range of redundant target effects, depending on its specific parameters and the nature of the evidence accumulation process (e.g., diffusion to bound or linear ballistic accumulation). More work is needed to test those assumptions and parameters.

Relation to the wider literature on serial versus parallel word recognition

Two related questions have fueled many studies of visual word recognition and reading: (1) *Can* multiple words be recognized in parallel? (2) In natural reading, *do* readers process multiple words in parallel? Both questions have been heavily debated. The redundant target effects explored in the present article are one way to investigate the first question, but not the second. The goal is to investigate the fundamental processing limits of visual word recognition – are readers capable of recognizing two words in parallel, when they are forced to try? Several distinct paradigms have been used to answer that question.

One such paradigm is the unspeeded dual-task paradigm that measures accuracy. It has provided evidence for a serial “bottleneck” in word recognition (White, Boynton, et al., 2019). In these experiments, participants were presented with two words at once. They were instructed either to respond to one pre-cued word (with focused attention) or to respond to both words in sequence (with divided attention). A key difference from the redundant target paradigm is that the two words must had to be judged independently,

rather than integrated to one decision. Also, the primary measure was accuracy rather than response time. These studies post-masked the words after a short interval calibrated to each individual's performance in the single-task condition. Thus, each participant was given just enough time to process one word, and the question was whether they can process two words with divided attention in that same amount of time. Standard parallel and serial processing models make different predictions for the magnitude of the drop in accuracy in the divided attention condition.

The results of several dual-task studies have been consistent with serial processing: the observer can recognize only one word per trial and must guess when asked about the other. That has been true for semantic categorization, lexical decision, and pronounceability judgments (Campbell et al., 2024; White et al., 2018, 2020; White, Palmer, et al., 2019). The large cost of dividing attention on accuracy in these experiments is consistent with a special case of the serial model in which only one word is processed. It rejects both the standard unlimited-capacity and the *fixed*-capacity parallel models. That measure alone cannot reject a more extreme limited-capacity parallel model. However, another result in these studies is a negative correlation between the two responses made within the same trial. The response to one stimulus was more likely to be correct when the response to the other stimulus was incorrect. This result is consistent with the standard serial model that processes just one of two words per trial, and rejects standard parallel models.

A related paradigm is “partially-valid cueing,” in which one of two stimulus locations is pre-cued as more likely to be task-relevant. This paradigm was applied to semantic judgments of words that were post-masked to limit the possibility of serially processing two words per trial. For the cued (more likely) location, accuracy was near 80% correct. For the uncued (less likely) location, accuracy was no better than chance (Johnson et al., 2022). This is again consistent with the version of the standard serial

model in which only one word can be processed (due to the time constraints that prevent switching). On the face of it, this result is inconsistent with any parallel model. To save the parallel model, one must assume a strategy of processing only one word at a time when one location is more likely to be relevant and time is limited; in other words, the parallel model becomes effectively serial.

The redundant target paradigm complements the dual-task and partially-valid cueing paradigms because the words do not have to be post-masked, and the observer does not need to make independent judgements about two words simultaneously. Moreover, the effects on response time can be compared to specific quantitative models. Altogether, the results are consistent in that they cannot reject the standard serial model.

Not all studies agree. Of particular interest is the flanker paradigm (Eriksen & Eriksen, 1974) that investigates whether judgments of a foveated target stimulus are affected by task-irrelevant flanking stimuli. This paradigm has been applied to words with several different tasks, demonstrating effects of the congruency between a target word and flanking words. Responses to a target are faster if the flankers are congruent with the target in terms of semantic or syntactic categories than if they are incongruent (Snell, Declerck, et al., 2018; Snell, Mathôt, et al., 2018; Snell, Meeter, et al., 2017; Snell & Grainger, 2018), but see (Broadbent & Gathercole, 1990). That is true even when the whole display is flashed for only 50 ms and then masked (Snell, 2024). Such flanker congruency effects have been interpreted as evidence for parallel processing of both the relevant target word and the irrelevant flanker words (Snell & Grainger, 2019).

In addition to these effects of higher order relations between words, the flanker paradigm provides evidence of sublexical orthographic effects. In particular, lexical decisions for a foveated target word are affected by non-word pairs of letters that flank it (Dare & Shillcock, 2013). Responses are faster when the flankers are bigrams that contain

similar letters as the target, compared to flankers that contain different letters (Grainger et al., 2014). That result supports a stage of parallel integration of orthographic information.

Importantly, this interactive parallel processing might primarily be deleterious. The integration of orthographic (letter-level) information across words can cause interference. In the lexical decision task, performance is generally worse when there are flankers of any type than when the target is presented alone (Snell, 2024; Snell, Declerck, et al., 2018). Further evidence for interference caused by flanking irrelevant words comes from the 'visual world' eye-movement paradigm (Ziaka et al., 2025). Moreover, "letter migration errors" illustrate a mixing of information: when asked to report the identity of one word, readers sometimes report a word that could be formed by combining the target's letters with letters from a neighboring word (Fischer-Baum et al., 2011; Mozer, 1983; Vandendaele et al., 2019). Facilitation (or at least, the absence of interference) may arise when a target word is flanked by a copy of itself (Snell & Grainger, 2018). Relatedly, in natural reading, fixation times on a target word are reduced if the next word over is an identical copy of it (Angele et al., 2013).

It might be tempting to explain the dual-task results reviewed above in terms of a flanker effect. Some flanker experiments highlight interference by comparing responses to one target word presented alone vs. responses to one target presented with other irrelevant words. In contrast, in the dual-task experiments, two words are always presented and the key manipulation is whether only one of the words is task-relevant, or whether both words are task-relevant. Thus, it is not trivial to explain the dual-task costs in terms of flanker effects.

In sum, flanker effects and related phenomena are consistent with some degree of parallel processing, starting at the sub-lexical orthographic level, which might in some cases impair word recognition but in other cases allow for parallel lexical, syntactic or

semantic processing. In the next section, we consider whether an interactive parallel processing model could explain our redundant target effects.

Coactivation and interactive parallel processing

In the literature on visual search and redundant targets, “coactivation” models assume that initially separate channels, which carry information about multiple stimuli, are combined before making a decision (Miller, 1982). We have not considered these models in detail, because they were designed to predict relatively large positive redundant target effects, in contrast to the negative effects we found.

Another class of models assume *interactive parallel processing*. One example of interactive processing is *crosstalk*: information about one stimulus affects the information that is represented about another stimulus. In essence, there is a failure of selectivity in the processing channels. In contrast, the specific models we tested in this article all assume “selective influence”: each stimulus is processed in a selective channel, such that the specific properties (e.g., letters or semantic category) of one stimulus do not affect the processing another stimulus.

Models with interactive processing might help explain our data and reconcile them with the flanker effect literature reviewed above. In fact, some authors have advanced a model of reading in which orthographic information is integrated across neighboring words into a “single channel” that then activates multiple word-level representations (Grainger et al., 2014, 2016; Snell, van Leipsig, et al., 2018). This idea has been proposed to explain the flanker effects reviewed above.

Such an interactive parallel model, with orthographic integration and interference, might explain the negative redundant target effects in our lexical and semantic tasks. Two target stimuli would both activate a single channel of letter or bigram detectors. These letter units would then activate ‘word units’ that contain those letters. Because the two targets were always different from each other, many potential word units would be

activated, and possibly inhibit each other. The activation of words that contain a mixture of letters from both targets would especially cause interference and slow responses compared to when a single target is presented alone. Thus, like some of our fixed-capacity parallel models, this interactive model might predict a negative redundant target effect on response times.

This model might also explain the *positive* redundant effects that arise when the two targets are identical words (Hasbrooke & Chiarello, 1998; Mohr et al., 1994; Mullin & Egeth, 1989). In the flanker task, responses to the target are not slowed by a single flanker immediately to the right of the target that is identical to the target (Snell & Grainger, 2018). Both results are consistent with pre-lexical pooling of letter identities, which can be beneficial when neighboring words are identical.

However, without more theoretical work, there are some difficulties in using an interactive parallel model to explain our redundant target effects. First without modifications, this model might predict a loss of *accuracy* on two-target trials. If something like ‘letter migrations’ were to occur (activations of words formed by mixing letters from the two targets; Mozer, 1983) a second target would introduce errors. In contrast, we always found that redundant targets increased accuracy, even when they slowed response times. The serial model can explain this benefit by having two chances to arrive at the correct decision (that a target was present).

Second, arguments for parallel processing on the basis of flanker effects highlight both a general cost of such parallel processing, due to orthographic interference, *and* an ability to extract syntactic or semantic information from multiple words at once (Snell, 2024; Snell & Grainger, 2019). In our experiments, that same type of parallel processing could have led to positive redundant target effects for the semantic task, which we did not find. It is possible, but not yet clear, that the negative effect might be explained by a

more complex mixture of orthographic interference and slowed but simultaneous semantic activations for both words (Snell, Vitu, et al., 2017).

The third difficulty for the parallel interactive model concerns the effect of the words' positions. Snell, Mathôt, et al. (2018) found strong flanker effects when the flankers were to the left and right of the target but not when they were above and below the target. In contrast, we found similar results for horizontal and vertical positions of the words. One explanation is that the parallel orthographic integration process occurs automatically across horizontally-arranged letter strings, but selective attention can more easily filter out irrelevant flankers that are above and below the target.

Thus, a parallel model with orthographic interference, as developed from flanker effects, might be able to explain negative redundant target effects, but there are some issues to resolve first. For now, we leave it as a possible alternative hypothesis that deserves further investigation.

Finally, let us briefly address one form of interactive processing that we do *not* believe is relevant to our findings: crowding. Crowding makes stimuli difficult to recognize when they are too close together, especially outside the fovea (Bouma, 1970; Pelli & Tillman, 2008). It applies to letters and many other types of stimuli. Crowding can be explained by an obligatory spatial pooling of the features of neighboring stimuli. The current experiments were designed to minimize this possibility by placing words on opposite sides of fixation.

Contingencies between stimuli

The effect of redundant targets can be modulated by contingencies among stimuli and between stimuli and responses (Mordkoff & Yantis, 1991). For example, one important contingency is the probability of observing a target at one location given the presence of either a target, distractor, or no stimulus at the other location. Participants can quickly

learn contingencies and adapt to them strategically. Contingency learning has been widely studied and perhaps the most relevant to our study is the large body of work on color-word contingency learning in the Stroop paradigm (Macleod, 2019; Schmidt, 2021).

Mordkoff and Yantis (1991) identified two kinds of contingency that can modulate the magnitude of the redundant target effect. One is the *nontarget-response contingency* that is not an issue in our experiments because it did not differ between set size conditions. We focus on the difference between two *interstimulus contingencies*: (i) the probability that one stimulus is a target given that the other location is blank. We denote this as $p(T_{L1} | B_{L2})$, where T means target, B means blank, and the subscripts L1, L2 denote the two locations. (ii) the probability that one stimulus is a target given that a target is present at the other location: $p(T_{L1} | T_{L2})$. If $p(T_{L1} | T_{L2}) > p(T_{L1} | B_{L2})$, then the contingency for the redundant target condition provides more information than the contingency for the single target condition. Mordkoff and Yantis (1991) demonstrated that such an imbalance in contingencies can increase the magnitude of the redundant target effect (by about 10 ms).

The two designs used in our experiments had different patterns of contingencies (compare Tables 2 and 5). In Experiments 1 and 2, $p(T_{L1} | B_{L2}) = 0.5$ and $p(T_{L1} | T_{L2}) = 0.67$. This favors a faster response for redundant targets compared to single targets. In contrast, in Experiments 3 and 4, the difference in interstimulus contingencies was in the opposite direction: $p(T_{L1} | B_{L2}) = 0.75$ and $p(T_{L1} | T_{L2}) = 0.33$. This favors a slower response for redundant targets compared to single targets.

Could these differences in interstimulus contingencies explain the negative redundant target effect we found for some tasks? We argue that they cannot, for two reasons: First, the interstimulus contingencies in Experiments 1 and 2 favored positive effects while in Experiments 3 and 4 they favored negative effects. Thus, they cannot have produced the negative effects found for the semantic task in *all experiments*. Second, the contingencies were the same for the color task and for the semantic task. Thus, they

cannot explain why the redundant target effect was reliably positive for the color task and negative for the semantic task. They might, however, explain why the redundant target effects for the lexical and semantic tasks were more robustly negative in Experiments 3–4 than in Experiments 1–2. In sum, while we have little doubt that contingencies can modulate the magnitude of redundant target effects, they cannot account for the negative effects we observed in the semantic tasks in all experiments.

Relation to the mimicry of parallel and serial models

Given that parallel and serial models can sometimes mimic one another (Algom et al., 2015; Townsend & Ashby, 1983), some have concluded that distinguishing such models is impossible. This is a misunderstanding (Townsend, 1990). But the issues are not simple, and we make three related points.

First, there is no contradiction between the analysis of model mimicry in the literature and the distinctive predictions of the specific models presented here. The mimicry analysis involves more general models, which generally make weaker predictions than specific models. Moreover, the “standard” models we presented here are still general in the sense that they assume no particular stochastic process or distribution. They are not equivalent and *can* be distinguished. Lastly, the prior models that mimic each other are specific in another sense: they are specific to variations of visual search tasks and do not apply to the larger variety of paradigms that have been used to investigate parallel processing of multiple words, such as dual tasks, partially valid cueing, and the flanker paradigm (reviewed above).

Second, we make progress by testing a specific hypothesis such as the standard serial model applied to word recognition. The strongest inferences in science come from rejecting a hypothesis. The challenge is that once a specific hypothesis is rejected, it is often possible to make additional assumptions to account for the data with a more complex version of the existing model. In fact, one way to think about model mimicry is

that it specifies an alternative hypothesis. Generating such alternative hypothesis is an important step in itself, but the never-ending sequence of alternative hypotheses can be disheartening.

Third, given the variety of different models of parallel and serial processing, we make progress by attacking the problem with a variety of tasks, stimuli, and measurements. The parallel vs. serial distinction for word recognition has been addressed in at least a half a dozen ways as reviewed above. In each case, a simple specific model (e.g. standard serial) might account for the data while the alternative specific model (e.g. standard unlimited-capacity parallel) requires some additional assumptions to remain viable. Parsimony favors the simpler model. But that by itself is not convincing. To reach consensus, one must repeat the argument over the entire set of relevant studies. If there is converging evidence over many domains consistent with a simpler model, then that model is preferred. In our opinion, the published data so far do not provide a convincing case for rejecting the standard serial model for written word recognition.

Relation to previous theory on response time and accuracy

Most previous work has followed one of two paths. One is to develop a pure response time theory that ignores errors (Townsend & Ashby, 1983; Townsend & Nozawa, 1995). This work has also been general in not assuming particular stochastic processes or response time distributions. The second path is to assume a specific stochastic process such as the diffusion process or the linear ballistic accumulator (both described in the Appendix; see Luce (1991) for an introduction). For example, Blurton et al. (2014) built on the diffusion process to model the redundant target effect. The strength of this path is the integrated treatment of response time and accuracy.

Here we sought to expand the general response time models to incorporate errors. The surprising result was that the standard, self-terminating, serial model with errors

showed a *negative* redundant target effect. This is not predicted by the corresponding pure response time model.

This work complements other recent effects to generalize pure response time models. In particular, Little et al. (2022) extended part of the theory of systems factorial technology to include errors. They examined the prediction of the double-factorial paradigm to distinguish parallel and serial processes. They showed that the previous analysis of exhaustive search models was general to conditions with errors. However, they did not find a similar general result for self-terminating search models. Instead, they examined two special cases and showed that the analysis for pure response time did generalize to those cases. This is important progress, but it remains to be determined if this method of distinguishing parallel and serial models holds for *all* standard, self-terminating models with errors.

In summary, a critical development is the creation of general theories of both response time and accuracy. We have developed such a theory for the redundant target paradigm.

Conclusion

This study makes two primary contributions: first, we generalized standard models of the redundant target effect, which yielded new predictions. These generalized models have several advantages: they are straightforward to implement and interpret; they do not assume any particular stochastic process or response time distribution, and they include errors while also modeling response time. The generalized standard, self-terminating serial model predicts that redundant targets can slow correct responses, even when they increase accuracy. In contrast, the standard, self-terminating, unlimited-capacity parallel model always predicts positive redundant target effects, even when allowing for errors. We also developed specific examples of standard, self-terminating,

fixed-capacity parallel models, some of which can predict negative redundant target effects.

Second, we presented experimental tests of these predictions for judgements of words. When the task required judgment of the letter colors, a positive redundant target effect rejected the standard self-terminating serial model. This is the result most commonly observed in the redundant target literature. But when the task required judgments of the words' meaning, a negative redundant target effect rejected the standard, self-terminating, unlimited-capacity parallel model and was instead consistent with either the standard self-terminating serial model or some variants of a fixed-capacity parallel model. This stands in contrast to most previous studies of word recognition that found positive redundant target effects and argued against the standard serial model.

In sum, this work demonstrates that negative redundant target effects do occur. They are consistent with some (but not all) fixed-capacity parallel models, or a standard serial model that can readily account for them. Thus, unlike for many other visual tasks, for word recognition the redundant target paradigm has not provided evidence against the standard serial model.

Acknowledgments: We are grateful to Nicole Oppenheimer for help with data collection. The National Eye Institute provided funding (grant R00 EY-029366 to ALW).

Conflict of interest statement: The authors have no conflicts of interest to declare.

Author contributions statement: Conceptualization, AW & JP; Investigation; AW, GS & JH; Methodology, AW & JP; Formal Analysis: AW; Theory development: JP and ZZ; Writing: AW, JP & GS.

Public data repository: All stimuli, data, and analysis code are available at: <https://osf.io/7kn9u/> (White, 2025).

REFERENCES

- Abrams, R. L., & Greenwald, A. G. (2000). Parts outweigh the whole (word) in unconscious analysis of meaning. *Psychological Science*, 11(2), 118–124. <https://doi.org/10.1111/1467-9280.00226>
- Algom, D., Eidels, A., Hawkins, R. X. D., Jefferson, B., & Townsend, J. T. (2015). Features of response times: Identification of cognitive mechanisms through mathematical modeling. In J. R. Busemeyer, Z. Wang, J. T. Townsend, & A. Eidels (Eds.), *The Oxford Handbook of Computational and Mathematical Psychology* (pp. 63–98). Oxford University Press.
- Angele, B., Tran, R., & Rayner, K. (2013). Parafoveal-foveal overlap can facilitate ongoing word identification during reading: Evidence from eye movements. *Journal of Experimental Psychology: Human Perception and Performance*, 39(2), 526–538. <https://doi.org/10.1037/a0029492>
- Benjamini, Y., & Hochberg, Y. (1995). Controlling the False Discovery Rate: A Practical and Powerful Approach to Multiple Testing. *Journal of the Royal Statistical Society. Series B (Methodological)*, 90(12), 289–300.
- Blurton, S. P., Greenlee, M. W., & Gondan, M. (2014). Multisensory processing of redundant information in go/no-go and choice responses. *Attention, Perception, and Psychophysics*, 76(4), 1212–1233. <https://doi.org/10.3758/s13414-014-0644-0>
- Bouma, H. (1970). Interaction effects in parafoveal letter recognition. *Nature*, 226, 177–178.
- Broadbent, D. E., & Gathercole, S. E. (1990). The processing of non-target words: Semantic or not? *Quarterly Journal of Experimental Psychology*, 42A(1), 3–37. <https://doi.org/10.1192/bjp.111.479.1009-a>
- Brown, S. D., & Heathcote, A. (2008). The simplest complete model of choice response time: Linear ballistic accumulation. *Cognitive Psychology*, 57(3), 153–178. <https://doi.org/10.1016/j.cogpsych.2007.12.002>
- Campbell, M., Oppenheimer, N., & White, A. L. (2024). Severe processing capacity limits for sub-lexical features of letter strings. *Attention, Perception, & Psychophysics*, 86, 643–652. <https://doi.org/10.3758/s13414-023-02830-1>
- Chatfield, C., & Theobald, C. M. (1973). Mixtures and Random Sums. *Journal of the Royal Statistical Society. Series D (The Statistician)*, 22(4), 281–287.
- Corballis, M. C. (2002). Hemispheric interactions in simple reaction time. *Neuropsychologia*, 40(4), 423–434. [https://doi.org/10.1016/S0028-3932\(01\)00097-5](https://doi.org/10.1016/S0028-3932(01)00097-5)
- Corbett, E. A., & Smith, P. L. (2020). A diffusion model analysis of target detection in near-threshold visual search. *Cognitive Psychology*, 120(October 2019), 101289. <https://doi.org/10.1016/j.cogpsych.2020.101289>
- Cox, G. E., & Criss, A. H. (2019). Parametric supplements to systems factorial analysis: Identifying interactive parallel processing using systems of accumulators. *Journal of Mathematical Psychology*, 92, 102247. <https://doi.org/10.1016/j.jmp.2019.01.004>

- Dare, N., & Shillcock, R. (2013). Serial and parallel processing in reading: Investigating the effects of parafoveal orthographic information on nonisolated word recognition. *The Quarterly Journal of Experimental Psychology*, 66(3), 487–504. <https://doi.org/10.1080/17470218.2012.703212>
- Donkin, C., Little, D. R., & Houpt, J. W. (2014). Assessing the speed-accuracy trade-off effect on the capacity of information processing. *Journal of Experimental Psychology: Human Perception and Performance*, 40(3), 1183–1202. <https://doi.org/10.1037/a0035947>
- Egeth, H. E., Folk, C. L., & Mullin, P. A. (1989). Spatial parallelism in the processing of lines, letters, and lexicality. In *Object perception: Structure and process* (pp. 19–52).
- Eidels, A., Donkin, C., Brown, S. D., & Heathcote, A. (2010). Converging measures of workload capacity. *Psychonomic Bulletin and Review*, 17(6), 763–771. <https://doi.org/10.3758/PBR.17.6.763>
- Eriksen, B. A., & Eriksen, C. W. (1974). Effects of noise letters upon the identification of a target letter in a nonsearch task. In *Perception & Psychophysics* (Vol. 16, Issue 1, pp. 143–149). <https://doi.org/10.3758/BF03203267>
- Eriksen, C. W., Goettl, B., St. James, J. D., & Fournier, L. R. (1989). Processing redundant signals: Coactivation, divided attention, or what? *Perception & Psychophysics*, 45(4), 356–370. <https://doi.org/10.3758/BF03204950>
- Fischer-Baum, S., Charny, J., & McCloskey, M. (2011). Both-edges representation of letter position in reading. *Psychonomic Bulletin and Review*, 18(6), 1083–1089. <https://doi.org/10.3758/s13423-011-0160-3>
- Fitousi, D. (2021). When two faces are not better than one: Serial limited-capacity processing with redundant-target faces. *Attention, Perception, & Psychophysics*, 3118–3134. <https://doi.org/10.3758/s13414-021-02335-9>
- Fournier, L. R., & Eriksen, C. W. (1990). Coactivation in the Perception of Redundant Targets. *Journal of Experimental Psychology: Human Perception and Performance*, 16(3), 538–550. <https://doi.org/10.1037/0096-1523.16.3.538>
- Gondan, M., Götze, C., & Greenlee, M. W. (2010). Redundancy gains in simple responses and go/no-go tasks. *Attention, Perception, and Psychophysics*, 72(6), 1692–1709. <https://doi.org/10.3758/APP.72.6.1692>
- Graham, N., Robson, J. G., & Nachmias, J. (1978). Grating summation in fovea and periphery. *Vision Research*, 18(7), 815–825. [https://doi.org/10.1016/0042-6989\(78\)90122-0](https://doi.org/10.1016/0042-6989(78)90122-0)
- Grainger, J., Dufau, S., & Ziegler, J. C. (2016). A Vision of Reading. *Trends in Cognitive Sciences*, 20(3), 171–179. <https://doi.org/10.1016/j.tics.2015.12.008>
- Grainger, J., Mathôt, S., & Vitu, F. (2014). Tests of a model of multi-word reading: Effects of parafoveal flanking letters on foveal word recognition. *Acta Psychologica*, 146(1), 35–40. <https://doi.org/10.1016/j.actpsy.2013.11.014>
- Grice, G. R., & Reed, J. M. (1992). What makes targets redundant? *Perception & Psychophysics*, 51(5), 437–442. <https://doi.org/10.3758/BF03211639>

- Hasbrooke, R. E., & Chiarello, C. (1998). Bihemispheric processing of redundant bilateral lexical information. *Neuropsychology*, 12(1), 78–94. <https://doi.org/10.1037/0894-4105.12.1.78>
- Hershenson, M. (1962). Reaction time as a measure of retroactive inhibition. *Journal of Experimental Psychology*, 63(3), 289–293. <https://doi.org/10.1037/h0055703>
- Johnson, M. L., Palmer, J., Moore, C. M., & Boynton, G. M. (2022). Evidence from partially valid cueing that words are processed serially. *Psychonomic Bulletin & Review*, 0123456789. <https://doi.org/10.3758/s13423-022-02230-w>
- Krekelberg, B. (2024). *Matlab Toolbox for Bayes Factor Analysis*. <https://doi.org/10.5281/ZENODO.13744717>
- Little, D. R., Yang, H., Eidels, A., & Townsend, J. T. (2022). Extending Systems Factorial Technology to Errored Responses. *Psychological Review*, 129(3), 484–512. <https://doi.org/10.1037/rev0000232>
- Luce, R. D. (1991). Response Times: Their Role in Inferring Elementary Mental Organization. In *Response Times: Their Role in Inferring Elementary Mental Organization*. Oxford University Press. <https://doi.org/10.1093/acprof:oso/9780195070019.001.0001>
- Macleod, C. M. (2019). Learning simple associations. *Canadian Psychology*, 60(1), 3–13. <https://doi.org/10.1037/cap0000170>
- Medler, D. A., & Binder, J. R. (2005). *MCWord: An on-Line orthographic database of the English language*. <http://www.neuro.mcw.edu/mcword/>
- Miller, J. (1982). Divided attention: evidence for coactivation with redundant signals. *Cognitive Psychology*, 14(2), 247–279. [https://doi.org/10.1016/0010-0285\(82\)90010-X](https://doi.org/10.1016/0010-0285(82)90010-X)
- Miller, J., & Lopes, A. (1988). Testing race models by estimating the smaller of two true mean or true median reaction times: An analysis of estimation bias. *Perception & Psychophysics*, 44(6), 513–524. <https://doi.org/10.3758/BF03207485>
- Mohr, B., Pulvermüller, F., Mittelstädt, K., & Rayman, J. (1996). Multiple simultaneous stimulus presentation facilitates lexical processing. *Neuropsychologia*, 34(10), 1003–1013. [https://doi.org/10.1016/0028-3932\(96\)00006-1](https://doi.org/10.1016/0028-3932(96)00006-1)
- Mohr, B., Pulvermüller, F., & Zaidel, E. (1994). Lexical decision after left, right and bilateral presentation of function words, content words and non-words: Evidence for interhemispheric interaction. *Neuropsychologia*, 32(1), 105–124. [https://doi.org/10.1016/0028-3932\(94\)90073-6](https://doi.org/10.1016/0028-3932(94)90073-6)
- Mordkoff, J. T., & Egeth, H. E. (1993). Response Time and Accuracy Revisited: Converging Support for the Interactive Race Model. *Journal of Experimental Psychology: Human Perception and Performance*, 19(5), 981–991. <https://doi.org/10.1037/0096-1523.19.5.981>
- Mordkoff, J. T., & Yantis, S. (1991). An Interactive Race Model of Divided Attention. *Journal of Experimental Psychology: Human Perception and Performance*, 17(2), 520–538.
- Mozer, M. C. (1983). Letter migration in word perception. *Journal of Experimental Psychology: Human Perception and Performance*, 9(4), 531–546. <https://doi.org/10.1037/0096-1523.9.4.531>

- Mullin, P. A., & Egeth, H. E. (1989). Capacity limitations in visual word processing. *Journal of Experimental Psychology: Human Perception and Performance*, 15(1), 111–123.
- Mullin, P. A., Egeth, H. E., & Mordkoff, J. T. (1988). Redundant-target detection and processing capacity: The problem of positional preferences. *Perception & Psychophysics*, 43(6), 607–610.
- Palmer, J., Huk, A. C., & Shadlen, M. N. (2005). The effect of stimulus strength on the speed and accuracy of a perceptual decision. *Journal of Vision*, 5(5), 376–404.
<https://doi.org/10.1167/5.5.1>
- Peirce, J., Gray, J. R., Simpson, S., MacAskill, M., Höchenberger, R., Sogo, H., Kastman, E., & Lindeløv, J. K. (2019). PsychoPy2: Experiments in behavior made easy. *Behavior Research Methods*, 51(1), 195–203. <https://doi.org/10.3758/s13428-018-01193-y>
- Pelli, D. G., & Tillman, K. A. (2008). The uncrowded window of object recognition. *Nature Neuroscience*, 11(10), 1129–1135. <https://doi.org/10.1038/nn.2187>
- Raab, D. H. (1962). Statistical facilitation of simple reaction times. *Transactions of the New York Academy of Sciences*, 24(5), 574–590. <https://doi.org/10.1111/j.2164-0947.1962.tb01433.x>
- Ridgway, N., Milders, M., & Sahraie, A. (2008). Redundant target effect and the processing of colour and luminance. *Experimental Brain Research*, 187(1), 153–160.
<https://doi.org/10.1007/s00221-008-1293-0>
- Robson, J. G., & Graham, N. (1981). Probability summation and regional variation in contrast sensitivity across the visual field. *Vision Research*, 21(3), 409–418.
[https://doi.org/10.1016/0042-6989\(81\)90169-3](https://doi.org/10.1016/0042-6989(81)90169-3)
- Rouder, J. N., Speckman, P. L., Sun, D., Morey, R. D., & Iverson, G. (2009). Bayesian t tests for accepting and rejecting the null hypothesis. *Psychonomic Bulletin & Review*, 16(2), 225–237. <https://doi.org/10.3758/PBR.16.2.225>
- Scharff, A., Palmer, J., & Moore, C. M. (2011). Extending the simultaneous-sequential paradigm to measure perceptual capacity for features and words. *Journal of Experimental Psychology: Human Perception and Performance*, 37(3), 813–833.
<https://doi.org/10.1037/a0021440>
- Schmidt, J. R. (2021). Incidental Learning of Simple Stimulus-Response Associations: A Review of Colour-Word Contingency Learning Research. *Annee Psychologique*, 121(2), 77–127.
<https://doi.org/10.3917/ANPSY1.212.0077>
- Schröter, H., Ulright, R., & Miller, J. (2007). Effects of redundant auditory stimuli on reaction time. *Psychonomic Bulletin & Review*, 14(1), 39–44.
- Schwarz, W. (2006). On the relationship between the redundant signals effect and temporal order judgments: Parametric data and a new model. *Journal of Experimental Psychology: Human Perception and Performance*, 32(3), 558–573. <https://doi.org/10.1037/0096-1523.32.3.558>
- Shaw, M. L. (1980). Identifying attentional and decision-making components in information processing. In R. Nickerson (Ed.), *Attention and Performance VIII* (pp. 277–296). Routledge. <https://doi.org/10.1017/CBO9781107415324.004>

- Shepherdson, P. V., & Miller, J. (2014). Redundancy gain in semantic categorisation. *Acta Psychologica*, 148, 96–106. <https://doi.org/10.1016/j.actpsy.2014.01.011>
- Snell, J. (2024). The reading brain extracts syntactic information from multiple words within 50 milliseconds. *Cognition*, 242(November 2023), 105664. <https://doi.org/10.1016/j.cognition.2023.105664>
- Snell, J., Declerck, M., & Grainger, J. (2018). Parallel semantic processing in reading revisited: effects of translation equivalents in bilingual readers. *Language, Cognition and Neuroscience*, 33(5), 563–574. <https://doi.org/10.1080/23273798.2017.1392583>
- Snell, J., & Grainger, J. (2018). Parallel word processing in the flanker paradigm has a rightward bias. *Attention, Perception, & Psychophysics*, 80(June), 1512–1519.
- Snell, J., & Grainger, J. (2019). Readers Are Parallel Processors. *Trends in Cognitive Sciences*, 23(7), 537–546. <https://doi.org/10.1016/j.tics.2019.04.006>
- Snell, J., Kempen, T. Van, & Olivers, C. N. L. (2022). Multi-res: An interface for improving reading without central vision. *Vision Research*, 201(April), 108129. <https://doi.org/10.1016/j.visres.2022.108129>
- Snell, J., Mathôt, S., Mirault, J., & Grainger, J. (2018). Parallel graded attention in reading: A pupillometric study. *Scientific Reports*, 8(1), 1–9. <https://doi.org/10.1038/s41598-018-22138-7>
- Snell, J., Meeter, M., & Grainger, J. (2017). Evidence for simultaneous syntactic processing of multiple words during reading. *PLoS ONE*, 12(3), 1–17. <https://doi.org/10.1371/journal.pone.0173720>
- Snell, J., van Leipsig, S., Grainger, J., & Meeter, M. (2018). OB1-reader: A model of word recognition and eye movements in text reading. *Psychological Review*, 125(6), 969–984. <https://doi.org/10.1037/rev0000119>
- Snell, J., Vitu, F., & Grainger, J. (2017). Integration of parafoveal orthographic information during foveal word reading: beyond the sub-lexical level? *Quarterly Journal of Experimental Psychology*, 70(10), 1984–1996. <https://doi.org/10.1080/17470218.2016.1217247>
- Taylor, M. M., Lindsay, P. H., & Forbes, S. M. (1967). Quantification of shared capacity processing in auditory and visual discrimination. In *Acta Psychologica* (Vol. 27, pp. 223–229). [https://doi.org/10.1016/0001-6918\(67\)99000-2](https://doi.org/10.1016/0001-6918(67)99000-2)
- Terry, A., Marley, A. A. J., Barnwal, A., Wagenmakers, E. J., Heathcote, A., & Brown, S. D. (2015). Generalising the drift rate distribution for linear ballistic accumulators. *Journal of Mathematical Psychology*, 68–69, 49–58. <https://doi.org/10.1016/j.jmp.2015.09.002>
- Thornton, T., & Gilden, D. L. (2001). Attentional Limitations in the Sensing of Motion Direction. *Cognitive Psychology*, 43(1), 23–52. <https://doi.org/10.1006/cogp.2001.0751>
- Townsend, J. T. (1990). Serial vs. parallel processing: Sometimes they look like Tweedledum and Tweedledee but they can (and should) be distinguished. *Psychological Science*, 1(1), 46–54. <https://doi.org/10.1038/sj/gt/3301593>

- Townsend, J. T., & Ashby, F. G. (1983). Stochastic Modeling of Elementary Psychological Processes. In *Cambridge University Press*. Cambridge University Press.
<https://doi.org/10.2307/1422636>
- Townsend, J. T., & Nozawa, G. (1995). Spatio-temporal Properties of Elementary Perception: An Investigation of Parallel, Serial, and Coactive Theories. In *Journal of Mathematical Psychology* (Vol. 39, Issue 4, pp. 321–359). <https://doi.org/10.1006/jmps.1995.1033>
- van der Heijden, A. H. C. (1975). Some evidence for a limited capacity parallel selfterminating process in simple visual search tasks. *Acta Psychologica*, 39(1), 21–41.
[https://doi.org/10.1016/0001-6918\(75\)90019-0](https://doi.org/10.1016/0001-6918(75)90019-0)
- van der Heijden, A. H. C., La Heij, W., & Boer, J. P. A. (1983). Parallel processing of redundant targets in simple visual search tasks. *Psychological Research*, 45(3), 235–254.
<https://doi.org/10.1007/BF00308704>
- Van Zandt, T., & Townsend, J. T. (1993). Self-terminating versus exhaustive processes in rapid visual and memory search: An evaluative review. *Perception & Psychophysics*, 53(5), 563–580. <https://doi.org/10.3758/BF03205204>
- Vandendaele, A., Snell, J., & Grainger, J. (2019). Letter migration errors reflect spatial pooling of orthographic information. *Attention, Perception, and Psychophysics*, 81(6), 2026–2036.
<https://doi.org/10.3758/s13414-019-01746-z>
- Verghese, P., & Stone, L. S. (1995). Combining speed information across space. *Vision Research*, 35(20), 2811–2823. [https://doi.org/10.1016/0042-6989\(95\)00038-2](https://doi.org/10.1016/0042-6989(95)00038-2)
- White, A. L. (2025, July 28). Repository of raw data and analysis code for White, Palmer et al. (2025). Retrieved from osf.io/7kn9u
- White, A. L., Boynton, G. M., & Yeatman, J. D. (2019). You Can't Recognize Two Words Simultaneously. *Trends in Cognitive Sciences*, 23(10), 812–814.
<https://doi.org/10.1016/j.tics.2019.07.001>
- White, A. L., Palmer, J., & Boynton, G. M. (2018). Evidence of serial processing in visual word recognition. *Psychological Science*, 29(7), 1062–1071.
<https://doi.org/10.1177/0956797617751898>
- White, A. L., Palmer, J., & Boynton, G. M. (2020). Visual word recognition: Evidence for a serial bottleneck in lexical access. *Attention, Perception, and Psychophysics*, 82, 2000–2017. <https://doi.org/10.3758/s13414-019-01916-z>
- White, A. L., Palmer, J., Boynton, G. M., & Yeatman, J. D. (2019). Parallel spatial channels converge at a bottleneck in anterior word-selective cortex. *Proceedings of the National Academy of Sciences*, 116(24), 10087–10096. <https://doi.org/10.1073/pnas.1822137116>
- Yeatman, J. D., & White, A. L. (2021). Reading: The Confluence of Vision and Language. *Annual Review of Vision Science*, 7(1), 487–517. <https://doi.org/10.1146/annurev-vision-093019-113509>
- Ziaka, L., Zelihic, D., McMurray, B., Baxelbaum, K., Simonsen, K., & Protopapas, A. (2025). Visual word recognition is impeded by adjacent words. *Journal of Memory and Language*, 143(April), 104641. <https://doi.org/10.1016/j.jml.2025.104641>

APPENDIX

This appendix describes three closely matched models of the redundant target effect: one serial and two parallel. These models start from the standard self-terminating search models of response time (Townsend & Nozawa, 1995) and add an account of accuracy. The main new result is that when there are errors, the standard self-terminating serial model predicts a slower response with two targets compared to one. This contrasts with the corresponding self-terminating serial model without errors that predicts no effect on response time. Another new result concerns the prediction of the standard self-terminating, unlimited-capacity, parallel model with errors. It predicts that the response time for two targets is faster than for one target. This is in accord with the corresponding parallel model without errors, although the effect is reduced with errors. Finally, for the standard fixed-capacity, parallel model there are no general predictions. Thus, among these landmark models, the redundant target paradigm can help distinguish serial and parallel processing.

Task Description

We focus on typical yes-no visual search tasks in which one or two stimuli are presented. A stimulus can either be a target t or a distractor d . When the task has one stimulus, the possible stimuli are t and d . When the task has two stimuli, the possible stimuli are tt , dd and td . Thus, the entire set of possible stimuli is $S = \{t, d, tt, dd, td\}$.

The task is to respond “yes” to the presence of any target, and respond “no” to the absence of any target.

There are four primary response measures of the redundant target task to be predicted that are subscripted by the stimulus condition: probability of a correct response p_s , the mean correct response time $\mu_{s,correct}$, the standard deviation of the correct response time $\sigma_{s,correct}$, and the mean incorrect response time $\mu_{s,incorrect}$ for $s \in S$. For example, for the single target condition t , these variables are denoted: p_t , $\mu_{t,correct}$, $\sigma_{t,correct}$ and $\mu_{t,incorrect}$.

Standard Self-terminating Serial Model with Errors

As with typical models of pure response time without errors, our serial model is based upon the selective influence of each stimulus on a separate process. In other words, there is one stimulus-specific component process for each stimulus. Each component process mediates the effect of one stimulus on both response time and accuracy.

For each possible stimulus $s \in \{t, d\}$ and associated component process, define a binary random variable for accuracy by $\mathbf{Z}_s$, which has a value of 1 if the decision is correct for that individual stimulus, and a value of 0 if incorrect. In addition, denote the continuous random variables for the component processing time for a correct decision for stimulus s by $\mathbf{D}_{s,correct}$, and for an incorrect decision by $\mathbf{D}_{s,incorrect}$. We emphasize

that these decisions are about a single stimulus and not about the response made to the set of stimuli.

As with similar models, assume that, besides the component processes, there are other “residual” processes that do not depend on the stimulus and do not affect accuracy, but contribute to the response time. The processing time from these residual processes is denoted by a continuous random variable $\mathbf{R}$ and it additively combines with the stimulus-specific component processes to yield the response time. We allow this residual processing time to depend on the specific response regardless of the stimulus. The random variable $\mathbf{R}_{yes}$ represents the residual processing time when the response is “yes” indicating the presence of a target, and $\mathbf{R}_{no}$ when the response is “no” indicating the absence of a target.

As with standard models of response times without errors, we assume a strong degree of independence, termed *context independence*, between component processes for different stimuli. Our definition has both an independent part (sometimes called stochastic independence) and an identical distribution part (sometimes called context invariance).

The component accuracy to one stimulus is assumed to be independent of other stimuli in the same stimulus condition. Consider a stimulus condition with two stimuli, denoted s_1 and s_2 , with $s_1, s_2 \in \{t, d\}$. Under this assumption, the random variable for component accuracy to stimulus s_1 , $\mathbf{Z}_{s_1}$, is independent of the random variable for

component accuracy to stimulus s_2 , $\mathbf{Z}_{s_2}$. Furthermore, when $s_1 = s_2$ (either both targets or both distractors), the component accuracies are identically distributed, that is, $\mathbf{Z}_{s_1}$ is identically distributed to $\mathbf{Z}_{s_2}$ (and identically distributed to either $\mathbf{Z}_t$ or $\mathbf{Z}_d$).

The component processing time for one stimulus is assumed to be independent of the component processing time of the other stimulus in the same stimulus condition. Specifically, for a stimulus condition with two stimuli s_1 and s_2 , with $s_1, s_2 \in \{t, d\}$, under this assumption the pairs of component processing times are independent. That is: $\mathbf{D}_{s_1,correct}$ is independent of $\mathbf{D}_{s_2,correct}$; $\mathbf{D}_{s_1,correct}$ is independent of $\mathbf{D}_{s_2,incorrect}$; $\mathbf{D}_{s_1,incorrect}$ is independent of $\mathbf{D}_{s_2,correct}$; and $\mathbf{D}_{s_1,incorrect}$ is independent of $\mathbf{D}_{s_2,incorrect}$. Furthermore, when $s_1 = s_2$ (either both targets or both distractors), the component processing times are identically distributed, that is, $\mathbf{D}_{s_1,correct}$ is identically distributed to $\mathbf{D}_{s_2,correct}$ (and identically distributed to either $\mathbf{D}_{t,correct}$ or $\mathbf{D}_{d,correct}$), and $\mathbf{D}_{s_1,incorrect}$ is identically distributed to $\mathbf{D}_{s_2,incorrect}$ (and identically distributed to either $\mathbf{D}_{t,incorrect}$ or $\mathbf{D}_{d,incorrect}$).

Additionally, the random variables $\mathbf{R}_{yes}$ and $\mathbf{R}_{no}$ are independent of other random variables. For example, when there is a target and the response is “yes”, $\mathbf{R}_{yes}$ is independent of $\mathbf{D}_{t,correct}$.

Context independence, including the accuracy and component processing assumptions, subsumes the more specific independence assumptions of independence from set size (unlimited capacity), independence from processing order (in the serial

model), and independence from the early completion of other processes (in the parallel model). For parallel models, we separate the unlimited-capacity assumption from context independence to allow consideration of limited capacity.

Even with context independence, the accuracy and component time within a single component process are not constrained and can be dependent. This allows them to be generated by a wide range of stochastic processes. For any $s \in \{t, d\}$, there is no restriction between $\mathbf{D}_{s,correct}$ given a correct decision, and $\mathbf{D}_{s,incorrect}$ given an incorrect decision. This is because correct and incorrect component processing times can differ, and are represented by separate variables. For example, the mean incorrect response time can be slower or faster than the mean correct response time.

Predictions of the Standard Self-terminating Serial Model with Errors

First consider stimulus conditions with either a single target t or two targets tt . Our goal is to describe the component processes of the two-stimulus conditions in terms of the component processes of a single stimulus.

For a **single target condition** t , the predicted probability of a correct response is defined as

$$p_t = P(\mathbf{Z}_t = 1). \quad (1)$$

The random variable for the correct response time to a target is,

$$\mathbf{T}_{t,correct} = \mathbf{D}_{t,correct} + \mathbf{R}_{yes}.$$

The expected response time for a correct response to a single target is then the sum of the expected values of that $\mathbf{D}_{t,correct}$ and $\mathbf{R}_{yes}$,

$$\mu_{t,correct} = E[\mathbf{D}_{t,correct}] + E[\mathbf{R}_{yes}]. \quad (2)$$

Relying on context independence, $\mathbf{D}_{t,correct}$ and $\mathbf{R}_{yes}$ are independent and therefore the variance of the response time of a correct response is the sum of the variances,

$$\sigma_{t,correct}^2 = \text{Var}[\mathbf{D}_{t,correct}] + \text{Var}[\mathbf{R}_{yes}]. \quad (3)$$

Similarly, the random variable for the incorrect response time to a target is,

$$\mathbf{T}_{t,incorrect} = \mathbf{D}_{t,incorrect} + \mathbf{R}_{no}$$

with expected value and variance as,

$$\mu_{t,incorrect} = E[\mathbf{D}_{t,incorrect}] + E[\mathbf{R}_{no}]$$

$$\sigma_{t,incorrect}^2 = \text{Var}[\mathbf{D}_{t,incorrect}] + \text{Var}[\mathbf{R}_{no}].$$

For the **two-target condition** tt , consider three mutually exclusive cases that describe the possible processing sequences of this serial model:

Case 1: One stimulus is processed first and is correct, and then processing is terminated (ignoring the other stimulus).

Case 2: One stimulus is processed first and is incorrect, and then the other stimulus is processed correctly.

Case 3: One stimulus is processed first and is incorrect, and then the other stimulus is processed and is also incorrect.

The probabilities and response times follow by case.

Case 1: The probability that Case 1 occurs is equal to the probability that a single stimulus is processed correctly, because of context independence,

$$p_{tt,Case1} = P(\mathbf{Z}_t = 1) = p_t.$$

The correct response time for Case 1 is the same as the response time for a correct response to a single target, also because of context independence,

$$\mathbf{T}_{tt,Case1} = \mathbf{D}_{t,correct} + \mathbf{R}_{yes}.$$

The expected correct response time for Case 1 is,

$$\mu_{tt,Case1} = E[\mathbf{D}_{t,correct}] + E[\mathbf{R}_{yes}].$$

Due to context independence, the variance of the correct response time for Case 1 is,

$$\sigma_{tt,Case1}^2 = \text{Var}[\mathbf{D}_{t,correct}] + \text{Var}[\mathbf{R}_{yes}].$$

Case 2: The probability that Case 2 occurs is equal to the probability that one stimulus is processed first and is incorrect, and that the other stimulus is processed correctly. Let t_{first} denote the target processed first, and t_{second} denote the target processed second. By context independence, the probability of Case 2 can be expressed in terms of single stimulus probabilities,

$$p_{tt,Case2} = P(\mathbf{Z}_{t_{first}} = 0 \text{ and } \mathbf{Z}_{t_{second}} = 1) = (1 - p_t)p_t.$$

The correct response time for Case 2, also relying on context independence, is the processing time for an incorrect response to one stimulus plus the processing time for a correct response to the other stimulus,

$$\mathbf{T}_{tt,Case2} = \mathbf{D}_{t,incorrect} + \mathbf{D}_{t,correct} + \mathbf{R}_{yes}.$$

The expected correct response time for Case 2 is,

$$\mu_{tt,Case2} = E[\mathbf{D}_{t,incorrect}] + E[\mathbf{D}_{t,correct}] + E[\mathbf{R}_{yes}].$$

The variance of the correct response time for Case 2, relying on context independence, is,

$$\sigma_{tt,Case2}^2 = \text{Var}[\mathbf{D}_{t,incorrect}] + \text{Var}[\mathbf{D}_{t,correct}] + \text{Var}[\mathbf{R}_{yes}].$$

Case 3: The probability that Case 3 occurs in the serial model is equal to the probability that both stimuli are processed incorrectly. By context independence,

$$p_{tt,Case3} = P(\mathbf{Z}_t = 0 \text{ and } \mathbf{Z}_t = 0) = (1 - p_t)^2.$$

The incorrect response time for Case 3 in the serial model is the processing time for an incorrect response to one stimulus plus the processing time for an incorrect response to the other stimulus,

$$\mathbf{T}_{tt,Case3} = 2 \mathbf{D}_{t,incorrect} + \mathbf{R}_{no}.$$

The expected incorrect response time for Case 3 is,

$$\mu_{tt,Case3} = 2E[\mathbf{D}_{t,incorrect}] + E[\mathbf{R}_{no}].$$

The variance of the incorrect response time for Case 3 is,

$$\sigma_{tt,Case3}^2 = 2 \text{Var}[\mathbf{D}_{t,incorrect}] + \text{Var}[\mathbf{R}_{no}].$$

In the two-target condition, a correct response is achieved in Case 1 *and* in Case 2, resulting in a mixture distribution (see (Chatfield & Theobald, 1973). The probability of a correct response is the probability of Case 1 plus the probability of Case 2, because they are mutually exclusive,

$$\begin{aligned}
p_{tt} &= p_{tt,Case1} + p_{tt,Case2} \\
&= p_t + (1 - p_t)p_t \\
&= 2p_t - p_t^2.
\end{aligned} \tag{4}$$

The expected response time for a correct response to a two-target condition, $\mu_{tt,correct}$, is the weighted average of the expected response times for Case 1 and Case 2. The weight for Case 1 is the proportion of correct responses for Case 1, i.e., the ratio of the probability of Case 1 to the probability of Case 1 or Case 2,

$$w_{Case1} = \frac{p_{tt,Case1}}{(p_{tt,Case1} + p_{tt,Case2})} = \frac{p_t}{(2p_t - p_t^2)} = \frac{1}{(2 - p_t)}.$$

Similarly,

$$w_{Case2} = \frac{p_{tt,Case2}}{(p_{tt,Case1} + p_{tt,Case2})} = \frac{p_t(1 - p_t)}{(2p_t - p_t^2)} = \frac{(1 - p_t)}{(2 - p_t)}.$$

Using these weights, the expected correct response time is,

$$\begin{aligned}
\mu_{tt,correct} &= w_{Case1}E[\mathbf{T}_{tt,Case1}] + w_{Case2}E[\mathbf{T}_{tt,Case2}] \\
&= w_{Case1}\mu_{tt,Case1} + w_{Case2}\mu_{tt,Case2} \\
&= \left(\frac{1}{2 - p_t}\right)E[\mathbf{D}_{t,correct} + \mathbf{R}_{yes}] + \left(\frac{1 - p_t}{2 - p_t}\right)E[\mathbf{D}_{t,incorrect} + \mathbf{D}_{t,correct} + \mathbf{R}_{yes}] \\
&= E[\mathbf{D}_{t,correct}] + \left(\frac{1 - p_t}{2 - p_t}\right)E[\mathbf{D}_{t,incorrect}] + E[\mathbf{R}_{yes}].
\end{aligned} \tag{5}$$

The variance of the correct response time for a mixture is the sum of two parts: the first part is the weighted averages of the variances of each case, and the second part is the variance due to the differences in the means of the cases. This is called the law of total variance (Chatman and Theobald, 1973), and yields,

$$\sigma_{tt,correct}^2 = w_{Case1} \sigma_{tt,Case1}^2 + w_{Case2} \sigma_{tt,Case2}^2 + \sigma_{tt,\Delta means}^2$$

where

$$\sigma_{tt,\Delta means}^2 = w_{Case1} \mu_{tt,Case1}^2 + w_{Case2} \mu_{tt,Case2}^2 - \mu_{tt,correct}^2.$$

The first part, the weighted averages of the variances of each case, can be expanded as

$$\begin{aligned} & w_{Case1} \sigma_{tt,Case1}^2 + w_{Case2} \sigma_{tt,Case2}^2 \\ &= \left(\frac{1}{2 - p_t} \right) (\text{Var}[\mathbf{D}_{t,correct}] + \text{Var}[\mathbf{R}_{yes}]) \\ &+ \left(\frac{1 - p_t}{2 - p_t} \right) (\text{Var}[\mathbf{D}_{t,correct}] + \text{Var}[\mathbf{D}_{t,incorrect}] + \text{Var}[\mathbf{R}_{yes}]) \\ &= \text{Var}[\mathbf{D}_{t,correct}] + \text{Var}[\mathbf{R}_{yes}] + \left(\frac{1 - p_t}{2 - p_t} \right) \text{Var}[\mathbf{D}_{t,incorrect}]. \end{aligned} \quad (6)$$

The second part, the variance due to the differences in the means of the cases, is always

non-negative, that is, $w_{Case1} \mu_{tt,Case1}^2 + w_{Case2} \mu_{tt,Case2}^2 - \mu_{tt,correct}^2 \geq 0$.

To obtain the expected incorrect response time, one can use the results for Case 3 because that is the only case that results in an incorrect response,

$$\mu_{tt,incorrect} = E[\mathbf{T}_{tt,Case3}] = 2E[\mathbf{D}_{t,incorrect}] + E[\mathbf{R}_{no}].$$

The variance of the incorrect response time is the same as that for Case 3,

$$\sigma_{tt,incorrect}^2 = 2 \text{Var}[\mathbf{D}_{t,incorrect}] + \text{Var}[\mathbf{R}_{no}].$$

Our main focus is on the difference between the correct response time for one target and the correct response time for two targets. The difference is constructed so that a faster response time to two targets results in a positive difference. Using Equations (2) and (5), the difference is,

$$\begin{aligned}
\mu_{t,correct} - \mu_{tt,correct} &= (E[\mathbf{D}_{t,correct}] + E[\mathbf{R}_{yes}]) \\
&\quad - \left(E[\mathbf{D}_{t,correct}] + \left(\frac{1-p_t}{2-p_t} \right) E[\mathbf{D}_{t,incorrect}] + E[\mathbf{R}_{yes}] \right) \\
&= - \left(\frac{1-p_t}{2-p_t} \right) E[\mathbf{D}_{t,incorrect}]. \tag{7}
\end{aligned}$$

This difference is less than or equal to zero. When there are no errors ($p_t = 1$), the correct response times are equal and difference equals zero. Thus, this serial model predicts, in the presence of errors, that a correct response time for two targets is slower than the correct response time for one target.

We also investigate the difference between the variance associated with a correct response time for one target and the variance associated with a correct response time for two targets. Using Equations (3) and (6), the difference is,

$$\begin{aligned}
\sigma_{t,correct}^2 - \sigma_{tt,correct}^2 &= \text{Var}[\mathbf{D}_{t,correct}] + \text{Var}[\mathbf{R}_{yes}] \\
&\quad - \left(\text{Var}[\mathbf{D}_{t,correct}] + \text{Var}[\mathbf{R}_{yes}] + \left(\frac{1-p_t}{2-p_t} \right) \text{Var}[\mathbf{D}_{t,incorrect}] \right) \\
&\quad - \sigma_{tt,\Delta means}^2 \\
&= - \left(\frac{1-p_t}{2-p_t} \right) \text{Var}[\mathbf{D}_{t,incorrect}] - \sigma_{tt,\Delta means}^2
\end{aligned}$$

which is less than or equal to zero. Thus, the serial model predicts, in the presence of errors, that the variance of a correct response time for two targets is larger than that for one target. See the final section of the Appendix for simulations of this effect on response time variability, and an explanation for why it is difficult to measure.

Next consider the difference between the probability of a correct response for two targets and the probability of a correct response for one target. The difference is constructed so that an increase in accuracy for two targets results in a positive difference. Using Equations (1) and (4), the difference is,

$$p_{tt} - p_t = (2 p_t - p_t^2) - p_t = p_t - p_t^2 \quad (8)$$

which is greater than or equal to zero. Thus, redundant targets improve accuracy.

For completeness, the other predictions of this serial model are given next.

For a **single distractor condition** d , by definition,

$$p_d = P(\mathbf{Z}_d = 1),$$

$$\mathbf{T}_{d,correct} = \mathbf{D}_{d,correct} + \mathbf{R}_{no}, \text{ and}$$

$$\mathbf{T}_{d,incorrect} = \mathbf{D}_{d,incorrect} + \mathbf{R}_{yes}.$$

The expected values and variances are,

$$\mu_{d,correct} = E[\mathbf{D}_{d,correct}] + E[\mathbf{R}_{no}],$$

$$\sigma_{d,correct}^2 = \text{Var}[\mathbf{D}_{d,correct}] + \text{Var}[\mathbf{R}_{no}],$$

$$\mu_{d,incorrect} = E[\mathbf{D}_{d,incorrect}] + E[\mathbf{R}_{yes}], \text{ and}$$

$$\sigma_{d,incorrect}^2 = \text{Var}[\mathbf{D}_{d,incorrect}] + \text{Var}[\mathbf{R}_{yes}].$$

For the **two-distractor condition** dd , there are three cases. The first case is when both distractors are processed correctly,

$$p_{dd,Case1} = P(\mathbf{Z}_d = 1 \text{ and } \mathbf{Z}_d = 1) = p_d^2, \text{ and}$$

$$\mathbf{T}_{dd,Case1} = 2\mathbf{D}_{d,correct} + \mathbf{R}_{no},$$

with

$$\mu_{dd,Case1} = 2 E[\mathbf{D}_{d,correct}] + E[\mathbf{R}_{no}],$$

$$\sigma_{dd,Case1}^2 = 2 \text{Var}[\mathbf{D}_{d,correct}] + \text{Var}[\mathbf{R}_{no}].$$

The second case is when one distractor is processed incorrectly, which terminates processing,

$$p_{dd,Case2} = P(\mathbf{Z}_d = 0) = 1 - p_d, \text{ and}$$

$$\mathbf{T}_{dd,Case2} = \mathbf{D}_{d,incorrect} + \mathbf{R}_{yes},$$

with

$$\mu_{dd,Case2} = E[\mathbf{D}_{d,incorrect}] + E[\mathbf{R}_{yes}],$$

$$\sigma_{dd,Case2}^2 = \text{Var}[\mathbf{D}_{d,incorrect}] + \text{Var}[\mathbf{R}_{yes}].$$

The third case is when one distractor is processed correctly but the second distractor is processed incorrectly,

$$p_{dd,Case3} = P(\mathbf{Z}_d = 1 \text{ and } \mathbf{Z}_d = 0) = p_d(1 - p_d), \text{ and}$$

$$\mathbf{T}_{dd,Case3} = \mathbf{D}_{d,correct} + \mathbf{D}_{d,incorrect} + \mathbf{R}_{yes},$$

with

$$\mu_{dd,Case3} = E[\mathbf{D}_{d,correct}] + E[\mathbf{D}_{d,incorrect}] + E[\mathbf{R}_{yes}],$$

$$\sigma_{dd,Case3}^2 = \text{Var}[\mathbf{D}_{d,correct}] + \text{Var}[\mathbf{D}_{d,incorrect}] + \text{Var}[\mathbf{R}_{yes}].$$

Only the first case yields a correct response, thus

$$p_{dd} = p_{dd,Case1} = p_d^2, \text{ and}$$

$$\mathbf{T}_{dd,correct} = 2\mathbf{D}_{d,correct} + \mathbf{R}_{no}.$$

The expected value and variance for a correct response is

$$\mu_{dd,correct} = 2 E[\mathbf{D}_{d,correct}] + E[\mathbf{R}_{no}],$$

$$\sigma_{dd,correct}^2 = 2 \text{Var}[\mathbf{D}_{d,correct}] + \text{Var}[\mathbf{R}_{no}].$$

The other two cases yield incorrect responses, yielding a mixture distribution. The weight for Case 2 is the fraction of incorrect responses due to Case 2 relative to all incorrect responses,

$$w_{Case2} = \frac{(1 - p_d)}{((1 - p_d) + p_d(1 - p_d))} = \frac{1}{1 + p_d}.$$

Similarly, the fraction of incorrect responses due to Case 3 relative to all incorrect responses is,

$$w_{Case3} = \frac{p_d(1 - p_d)}{((1 - p_d) + p_d(1 - p_d))} = \frac{p_d}{1 + p_d}.$$

The weighted combination of Cases 2 and 3 yields the expected incorrect response time,

$$\begin{aligned} \mu_{dd,incorrect} &= w_{Case2} E[\mathbf{T}_{dd,Case2}] + w_{Case3} E[\mathbf{T}_{dd,Case3}] \\ &= \left(\frac{1}{1 + p_d} \right) (E[\mathbf{D}_{d,incorrect}] + E[\mathbf{R}_{yes}]) \\ &\quad + \left(\frac{p_d}{1 + p_d} \right) (E[\mathbf{D}_{d,correct}] + E[\mathbf{D}_{d,incorrect}] + E[\mathbf{R}_{yes}]) \\ &= E[\mathbf{D}_{d,incorrect}] + \left(\frac{p_d}{1 + p_d} \right) E[\mathbf{D}_{d,correct}] + E[\mathbf{R}_{yes}]. \end{aligned}$$

The variance of the incorrect response time for the mixture of Cases 2 and 3 is,

$$\begin{aligned} \sigma_{dd,incorrect}^2 &= (w_{Case2} \sigma_{dd,Case2}^2 + w_{Case3} \sigma_{dd,Case3}^2) \\ &\quad + (w_{Case2} \mu_{dd,Case2}^2 + w_{Case3} \mu_{dd,Case3}^2 - \mu_{dd,incorrect}^2). \end{aligned}$$

For the **one target and one distractor condition** td , there are six mutually exclusive cases. The cases are distinguished by whether the target is processed first

(Cases 1, 2, and 3), or the distractor is processed first (Cases 4, 5, and 6). Which stimuli is processed first is considered to be random (typically with probability of 0.5). The first case is that the target is processed first and is correct,

$$p_{td,Case1} = p_t,$$

$$T_{td,Case1} = D_{t,correct} + R_{yes},$$

with

$$\mu_{td,Case1} = E[D_{t,correct}] + E[R_{yes}],$$

$$\sigma_{td,Case1}^2 = \text{Var}[D_{t,correct}] + \text{Var}[R_{yes}].$$

The second case is that the target is processed first incorrectly and then the distractor is processed correctly,

$$p_{td,Case2} = (1 - p_t)p_d$$

$$T_{td,Case2} = D_{t,incorrect} + D_{d,correct} + R_{no},$$

with

$$\mu_{td,Case2} = E[D_{t,incorrect}] + E[D_{d,correct}] + E[R_{no}],$$

$$\sigma_{td,Case2}^2 = \text{Var}[D_{t,incorrect}] + \text{Var}[D_{d,correct}] + \text{Var}[R_{no}].$$

The third case is that the target is processed first incorrectly and then the distractor is processed incorrectly,

$$p_{td,Case3} = (1 - p_t)(1 - p_d),$$

$$T_{td,Case3} = D_{t,incorrect} + D_{d,incorrect} + R_{yes},$$

with

$$\mu_{td,Case3} = E[D_{t,incorrect}] + E[D_{d,incorrect}] + E[R_{yes}],$$

$$\sigma_{td,Case3}^2 = \text{Var}[\mathbf{D}_{t,incorrect}] + \text{Var}[\mathbf{D}_{d,incorrect}] + \text{Var}[\mathbf{R}_{yes}].$$

The fourth case is that the distractor is processed first incorrectly, which terminates processing,

$$p_{td,Case4} = (1 - p_d),$$

$$\mathbf{T}_{td,Case4} = \mathbf{D}_{d,incorrect} + \mathbf{R}_{yes},$$

with

$$\mu_{td,Case4} = \text{E}[\mathbf{D}_{d,incorrect}] + \text{E}[\mathbf{R}_{yes}],$$

$$\sigma_{td,Case4}^2 = \text{Var}[\mathbf{D}_{d,incorrect}] + \text{Var}[\mathbf{R}_{yes}].$$

The fifth case is that the distractor is processed first correctly and then the target is processed correctly,

$$p_{td,Case5} = p_d p_t,$$

$$\mathbf{T}_{td,Case5} = \mathbf{D}_{d,correct} + \mathbf{D}_{t,correct} + \mathbf{R}_{yes},$$

with

$$\mu_{td,Case5} = \text{E}[\mathbf{D}_{d,correct}] + \text{E}[\mathbf{D}_{t,correct}] + \text{E}[\mathbf{R}_{yes}],$$

$$\sigma_{td,Case5}^2 = \text{Var}[\mathbf{D}_{d,correct}] + \text{Var}[\mathbf{D}_{t,correct}] + \text{Var}[\mathbf{R}_{yes}].$$

The sixth case is that the distractor is processed first correctly and then the target is processed incorrectly,

$$p_{td,Case6} = p_d (1 - p_t),$$

$$\mathbf{T}_{td,Case6} = \mathbf{D}_{d,correct} + \mathbf{D}_{t,incorrect} + \mathbf{R}_{no},$$

with

$$\mu_{td,Case6} = \text{E}[\mathbf{D}_{d,correct}] + \text{E}[\mathbf{D}_{t,incorrect}] + \text{E}[\mathbf{R}_{no}],$$

$$\sigma_{td,Case6}^2 = \text{Var}[\mathbf{D}_{d,correct}] + \text{Var}[\mathbf{D}_{t,incorrect}] + \text{Var}[\mathbf{R}_{no}].$$

The probability of a correct response is achieved through the mutually exclusive cases 1, 3, 4 and 5. Specifically, it is the probability that the target is processed first (denoted p_{tfirst}) and results in Case 1 or Case 3, plus the probability that the distractor is processed first ($1 - p_{\text{tfirst}}$) and results in Case 4 or Case 5,

$$\begin{aligned} p_{td} &= p_{\text{tfirst}}(p_t + (1 - p_t)(1 - p_d)) + (1 - p_{\text{tfirst}})((1 - p_d) + p_t p_d) \\ &= 1 - p_d + p_t p_d. \end{aligned}$$

For the four cases that contribute to a correct response, the weights are

$$\begin{aligned} w_1 &= \frac{p_{\text{tfirst}} p_{td,Case1}}{p_{td}} = \frac{p_{\text{tfirst}} p_t}{(1 - p_d + p_t p_d)}, \\ w_3 &= \frac{p_{\text{tfirst}} p_{td,Case3}}{p_{td}} = \frac{p_{\text{tfirst}} (1 - p_t)(1 - p_d)}{(1 - p_d + p_t p_d)}, \\ w_4 &= \frac{(1 - p_{\text{tfirst}}) p_{td,Case4}}{p_{td}} = \frac{(1 - p_{\text{tfirst}}) (1 - p_d)}{(1 - p_d + p_t p_d)}, \\ w_5 &= \frac{(1 - p_{\text{tfirst}}) p_{td,Case5}}{p_{td}} = \frac{(1 - p_{\text{tfirst}}) p_t p_d}{(1 - p_d + p_t p_d)}. \end{aligned}$$

Using these weights, the expected correct response time is,

$$\begin{aligned} \mu_{td,correct} &= w_1 E[\mathbf{T}_{td,Case1}] + w_3 E[\mathbf{T}_{td,Case3}] + w_4 E[\mathbf{T}_{td,Case4}] + w_5 E[\mathbf{T}_{td,Case5}] \\ &= \frac{p_{\text{tfirst}} p_t}{(1 - p_d + p_t p_d)} (E[\mathbf{D}_{t,correct}] + E[\mathbf{R}_{yes}]) \\ &\quad + \frac{p_{\text{tfirst}} (1 - p_t)(1 - p_d)}{(1 - p_d + p_t p_d)} (E[\mathbf{D}_{t,incorrect}] + E[\mathbf{D}_{d,incorrect}] + E[\mathbf{R}_{yes}]) \\ &\quad + \frac{(1 - p_{\text{tfirst}}) (1 - p_d)}{(1 - p_d + p_t p_d)} (E[\mathbf{D}_{d,incorrect}] + E[\mathbf{R}_{yes}]) \\ &\quad + \frac{(1 - p_{\text{tfirst}}) p_t p_d}{(1 - p_d + p_t p_d)} (E[\mathbf{D}_{d,correct}] + E[\mathbf{D}_{t,correct}] + E[\mathbf{R}_{yes}]) \end{aligned}$$

$$\begin{aligned}
&= \frac{p_t p_d + p_{\text{tfirst}} p_t (1 - p_d)}{(1 - p_d + p_t p_d)} E[\mathbf{D}_{t,\text{correct}}] + \frac{p_{\text{tfirst}} (1 - p_t)(1 - p_d)}{(1 - p_d + p_t p_d)} E[\mathbf{D}_{t,\text{incorrect}}] \\
&\quad + \frac{(1 - p_{\text{tfirst}}) p_t p_d}{(1 - p_d + p_t p_d)} E[\mathbf{D}_{d,\text{correct}}] \\
&\quad + \frac{(1 - p_d) + p_{\text{tfirst}} (-p_t + p_t p_d)}{(1 - p_d + p_t p_d)} E[\mathbf{D}_{d,\text{incorrect}}] + E[\mathbf{R}_{\text{yes}}].
\end{aligned}$$

The variance of the correct response time as a mixture of Cases 1, 3, 4, and 5 is,

$$\begin{aligned}
\sigma_{td,\text{correct}}^2 &= (w_1 \sigma_{td,\text{Case1}}^2 + w_3 \sigma_{td,\text{Case3}}^2 + w_4 \sigma_{td,\text{Case4}}^2 + w_5 \sigma_{td,\text{Case5}}^2) \\
&\quad + (w_1 \mu_{td,\text{Case1}}^2 + w_3 \mu_{td,\text{Case3}}^2 + w_4 \mu_{td,\text{Case4}}^2 + w_5 \mu_{td,\text{Case5}}^2 - \mu_{td,\text{correct}}^2).
\end{aligned}$$

Similarly, the expected incorrect response time is due to Cases 2 and 6. The weights are

$$\begin{aligned}
w_2 &= \frac{p_{\text{tfirst}} p_{td,\text{Case2}}}{(1 - p_{td})} = \frac{p_{\text{tfirst}} (1 - p_t) p_d}{p_d (1 - p_t)} = p_{\text{tfirst}}, \\
w_6 &= \frac{(1 - p_{\text{tfirst}}) p_{td,\text{Case6}}}{(1 - p_{td})} = \frac{(1 - p_{\text{tfirst}}) p_d (1 - p_t)}{p_d (1 - p_t)} = (1 - p_{\text{tfirst}}).
\end{aligned}$$

Using these weights, the expected incorrect response time is,

$$\begin{aligned}
\mu_{td,\text{incorrect}} &= w_2 E[\mathbf{T}_{td,\text{Case2}}] + w_6 E[\mathbf{T}_{td,\text{Case6}}] \\
&= p_{\text{tfirst}} (E[\mathbf{D}_{t,\text{incorrect}}] + E[\mathbf{D}_{d,\text{correct}}] + E[\mathbf{R}_{no}]) \\
&\quad + (1 - p_{\text{tfirst}}) (E[\mathbf{D}_{d,\text{correct}}] + E[\mathbf{D}_{t,\text{incorrect}}] + E[\mathbf{R}_{no}]) \\
&= E[\mathbf{D}_{t,\text{incorrect}}] + E[\mathbf{D}_{d,\text{correct}}] + E[\mathbf{R}_{no}].
\end{aligned}$$

The variance of the incorrect response time as a mixture of Cases 2 and 6 is,

$$\sigma_{td,\text{incorrect}}^2 = (w_2 \sigma_{td,\text{Case2}}^2 + w_6 \sigma_{td,\text{Case6}}^2) + (w_2 \mu_{td,\text{Case2}}^2 + w_6 \mu_{td,\text{Case6}}^2 - \mu_{td,\text{incorrect}}^2).$$

Standard Self-terminating, Unlimited-capacity, Parallel Model with Errors

We generalize the standard self-terminating, unlimited-capacity, parallel model without errors to include errors, and ask whether the introduction of errors changes its

predictions. In the literature, the parallel model without errors predicts that the correct response time for two targets is faster than for one target. We examine this prediction when including errors in our generalized parallel model. The task description discussed in the first section is the same, and the notation is the same as introduced for the Standard Self-terminating Serial Model with Errors. The difference from the serial model is that, instead of component processes being executed sequentially, in the parallel model the component processes are executed in parallel. As with the serial model, context independence is assumed with a strong degree of independence between the component processes for different stimuli.

Predictions of the Standard Self-terminating, Unlimited-capacity, Parallel Model with Errors

Consider stimulus conditions with either a single target t or two targets tt . As before, our goal is to describe the processes of the two-stimulus conditions in terms of the single-stimulus component processes. For a single target, the parallel model has the same definition and corresponding prediction as the serial model.

For a **single target condition** t , the predicted probability of a correct response is

$$p_t = P(\mathbf{Z}_t = 1). \quad (9)$$

The random variable for the correct response time to a target is,

$$\mathbf{T}_{t,correct} = \mathbf{D}_{t,correct} + \mathbf{R}_{yes}.$$

The expected response time for a correct response to a single target is then the sum of the expected values of that $\mathbf{D}_{t,correct}$ and $\mathbf{R}_{yes}$,

$$\mu_{t,correct} = E[\mathbf{D}_{t,correct}] + E[\mathbf{R}_{yes}]. \quad (10)$$

The variance of the response time of a correct response is,

$$\sigma_{t,correct}^2 = \text{Var}[\mathbf{D}_{t,correct}] + \text{Var}[\mathbf{R}_{yes}]. \quad (11)$$

Similarly, the random variable for the incorrect response time to a target is,

$$\mathbf{T}_{t,incorrect} = \mathbf{D}_{t,incorrect} + \mathbf{R}_{no}$$

with expected value as,

$$\mu_{t,incorrect} = E[\mathbf{D}_{t,incorrect}] + E[\mathbf{R}_{no}].$$

For the **two-target condition** tt , consider three mutually exclusive cases that describe the possible processing sequences of the parallel model:

Case 1: Each stimulus is processed correctly in parallel. The response is reported as soon as a target is detected whichever process is completed first, and then processing is terminated, even though the other stimulus is partially processed.

Case 2: One stimulus is processed correctly, and the other stimulus is processed incorrectly. The response is reported when the stimulus that is processed correctly completes processing, regardless of whether the other stimulus has completed processing or is partially processed.

Case 3: Each stimulus is processed incorrectly. The response is reported after both stimuli have been processed.

The probabilities and response times follow by case.

Case 1: The probability that Case 1 occurs is equal to the probability that both stimuli are processed correctly. Using the independence of component accuracy of one stimulus (t_1) to the other (t_2), the probability that Case 1 occurs is,

$$p_{tt,Case1} = P(\mathbf{Z}_{t_1} = 1, \mathbf{Z}_{t_2} = 1) = p_t^2.$$

The correct response time for Case 1 is the processing time for the stimulus that was completed first plus the residual time, also because of context independence,

$$\mathbf{T}_{tt,Case1} = \min \{\mathbf{D}_{t_1,correct}, \mathbf{D}_{t_2,correct}\} + \mathbf{R}_{yes}.$$

The expected correct response time for Case 1 is,

$$\mu_{tt,Case1} = E[\min \{\mathbf{D}_{t_1,correct}, \mathbf{D}_{t_2,correct}\}] + E[\mathbf{R}_{yes}]$$

and the variance is

$$\sigma^2_{tt,Case1} = \text{Var}[\min\{\mathbf{D}_{t_1,correct}, \mathbf{D}_{t_2,correct}\}] + \text{Var}[\mathbf{R}_{yes}]. \quad (12)$$

Case 2: The probability that Case 2 occurs is equal to the probability that one stimulus is processed correctly and that the other stimulus is processed incorrectly. By context independence, the probability that Case 2 occurs can be expressed in terms of single stimulus probabilities,

$$p_{tt,Case2} = P(\mathbf{Z}_{t_1} = 0 \text{ and } \mathbf{Z}_{t_2} = 1, \text{ or } \mathbf{Z}_{t_1} = 1 \text{ and } \mathbf{Z}_{t_2} = 0) = 2(1 - p_t)p_t.$$

The correct response time for Case 2 is the time that one stimulus is processed correctly, regardless of whether the other stimulus has completed processing or is partially processed. The response has to wait until the correct response has completed processing, no matter whether the incorrect component processing time is greater or less than the correct component processing time. Again, by context independence, the correct response time for Case 2 is,

$$T_{tt,Case2} = D_{t,correct} + R_{yes}.$$

The expected correct response time for Case 2 is,

$$\mu_{tt,Case2} = E[D_{t,correct}] + E[R_{yes}]$$

and the variance is

$$\sigma_{tt,Case2}^2 = \text{Var}[D_{t,correct}] + \text{Var}[R_{yes}]. \quad (13)$$

Case 3: The probability that Case 3 occurs in the parallel model is equal to the probability that both stimuli are processed incorrectly. By context independence,

$$p_{tt,Case3} = P(Z_{t_1} = 0 \text{ and } Z_{t_2} = 0) = (1 - p_t)^2.$$

The incorrect response time for Case 3 in the parallel model is the longest processing time for an incorrect response plus the residual time,

$$T_{tt,Case3} = \max \{D_{t_1,incorrect}, D_{t_2,incorrect}\} + R_{no}.$$

The expected incorrect response time for Case 3 is,

$$\mu_{tt,Case3} = E[\max \{D_{t_1,incorrect}, D_{t_2,incorrect}\}] + E[R_{no}].$$

In the two-target condition, a correct response is achieved in Case 1 *and* in Case 2, resulting in a mixture distribution. The probability of a correct response is the probability of Case 1 plus the probability of Case 2, because they are mutually exclusive,

$$\begin{aligned}
 p_{tt} &= p_{tt,Case1} + p_{tt,Case2} \\
 &= p_t p_t + 2(1 - p_t)p_t \\
 &= 2 p_t - p_t^2.
 \end{aligned} \tag{14}$$

The expected response time for a correct response to a two-target condition, $\mu_{tt,correct}$, is the weighted average of the expected response times for Case 1 and Case 2. The weight for Case 1 is the proportion of correct responses for Case 1, i.e., the ratio of the probability of Case 1 to the probability of Case 1 or Case 2,

$$w_{Case1} = \frac{p_{tt,Case1}}{(p_{tt,Case1} + p_{tt,Case2})} = \frac{p_t p_t}{(2p_t - p_t^2)} = \frac{p_t}{(2 - p_t)}.$$

Similarly,

$$w_{Case2} = \frac{p_{tt,Case2}}{(p_{tt,Case1} + p_{tt,Case2})} = \frac{2p_t(1 - p_t)}{(2p_t - p_t^2)} = \frac{2(1 - p_t)}{(2 - p_t)}.$$

Using these weights, the expected correct response time is,

$$\begin{aligned}
 \mu_{tt,correct} &= w_{Case1} E[\mathbf{T}_{tt,Case1}] + w_{Case2} E[\mathbf{T}_{tt,Case2}] \\
 &= w_{Case1} \mu_{tt,Case1} + w_{Case2} \mu_{tt,Case2} \\
 &= \left(\frac{p_t}{2 - p_t} \right) (E[\min\{\mathbf{D}_{t_1,correct}, \mathbf{D}_{t_2,correct}\}] + E[\mathbf{R}_{yes}]) \\
 &\quad + \left(\frac{2(1 - p_t)}{(2 - p_t)} \right) (E[\mathbf{D}_{t,correct}] + E[\mathbf{R}_{yes}])
 \end{aligned}$$

$$= \left(\frac{p_t}{2 - p_t} \right) E[\min \{ \mathbf{D}_{t_1, correct}, \mathbf{D}_{t_2, correct} \}] + \left(\frac{2(1 - p_t)}{(2 - p_t)} \right) E[\mathbf{D}_{t, correct}] + E[\mathbf{R}_{yes}]. \quad (15)$$

The variance for the correct response time is

$$\begin{aligned} \sigma_{tt, correct}^2 &= w_{Case1} (Var[\min \{ \mathbf{D}_{t_1, correct}, \mathbf{D}_{t_2, correct} \}] + Var[\mathbf{R}_{yes}]) \\ &\quad + w_{Case2} (Var[\mathbf{D}_{t, correct}] + Var[\mathbf{R}_{yes}]) + \sigma_{tt, \Delta means}^2 \end{aligned} \quad (16)$$

where

$$\sigma_{tt, \Delta means}^2 = w_{Case1} \mu_{tt, Case1}^2 + w_{Case2} \mu_{tt, Case2}^2 - \mu_{tt, correct}^2.$$

The incorrect response time is given by Case 3 alone. The expected incorrect response time is,

$$\mu_{tt, incorrect} = E[\max \{ \mathbf{D}_{t_1, incorrect}, \mathbf{D}_{t_2, incorrect} \}] + E[\mathbf{R}_{no}].$$

Our main focus is on the difference between the expected correct response time for one target and the expected correct response time for two targets. The difference is constructed so that a faster response time to two targets results in a positive difference.

Using Equations (10) and (15), the difference is,

$$\begin{aligned} \mu_{t, correct} - \mu_{tt, correct} &= (E[\mathbf{D}_{t, correct}] + E[\mathbf{R}_{yes}]) \\ &\quad - \left(\left(\frac{p_t}{2 - p_t} \right) E[\min \{ \mathbf{D}_{t_1, correct}, \mathbf{D}_{t_2, correct} \}] + \left(\frac{2(1 - p_t)}{(2 - p_t)} \right) E[\mathbf{D}_{t, correct}] \right. \\ &\quad \left. + E[\mathbf{R}_{yes}] \right) \\ &= \left(\frac{2 - p_t - 2(1 - p_t)}{2 - p_t} \right) E[\mathbf{D}_{t, correct}] - \left(\frac{p_t}{2 - p_t} \right) E[\min \{ \mathbf{D}_{t_1, correct}, \mathbf{D}_{t_2, correct} \}] \end{aligned}$$

$$= \left(\frac{p_t}{2 - p_t} \right) (E[\mathbf{D}_{t,correct}] - E[\min\{\mathbf{D}_{t_1,correct}, \mathbf{D}_{t_2,correct}\}]). \quad (17)$$

The difference is positive, since $E[\mathbf{D}_{t,correct}] \geq E[\min\{\mathbf{D}_{t_1,correct}, \mathbf{D}_{t_2,correct}\}]$. Thus, this parallel model predicts that a correct response time for two targets is faster than the correct response time for one target.

We also investigate the difference between the variance associated with a correct response time for one target and the variance associated with a correct response time for two targets. Using Equations (11) and (16), the difference is,

$$\begin{aligned} \sigma_{t,correct}^2 - \sigma_{tt,correct}^2 &= \text{Var}[\mathbf{D}_{t,correct}] + \text{Var}[\mathbf{R}_{yes}] \\ &\quad - (w_{Case1}(\text{Var}[\min\{\mathbf{D}_{t_1,correct}, \mathbf{D}_{t_2,correct}\}] + \text{Var}[\mathbf{R}_{yes}]) \\ &\quad + w_{Case2}(\text{Var}[\mathbf{D}_{t,correct}] + \text{Var}[\mathbf{R}_{yes}]) + \sigma_{tt,\Delta means}^2). \end{aligned}$$

Notice that $w_{Case1} + w_{Case2} = 1$, and $\text{Var}[\mathbf{D}_{t,correct}] - w_{Case2}\text{Var}[\mathbf{D}_{t,correct}] = w_{Case1}\text{Var}[\mathbf{D}_{t,correct}]$, so this reduces to

$$\begin{aligned} \sigma_{t,correct}^2 - \sigma_{tt,correct}^2 &= w_{Case1}(\text{Var}[\mathbf{D}_{t,correct}] - \text{Var}[\min\{\mathbf{D}_{t_1,correct}, \mathbf{D}_{t_2,correct}\}]) - \sigma_{tt,\Delta means}^2. \end{aligned}$$

The difference of the component variances in the first two terms is positive while the value of $-\sigma_{tt,\Delta means}^2$ is negative. Our evaluation of specific models finds the difference in variance is positive but we do not have a proof that this is always the case. Thus, it is common that the one target condition has a larger variance than the two target condition, but it might not be universal.

Next consider the difference between the probability of a correct response for two targets and the probability of a correct response for one target. The difference is constructed so that an increase in accuracy for two targets results in a positive difference. Using Equations (9) and (14), the difference is,

$$(2 p_t - p_t^2) - p_t = p_t - p_t^2 \quad (18)$$

which is greater than or equal to zero because $p_t \geq p_t^2$. Thus, redundant targets improve accuracy.

For completeness, the other predictions of this parallel model are given next.

For a **single distractor condition** d , by definition,

$$p_d = P(\mathbf{Z}_d = 1),$$

$$\mathbf{T}_{d,correct} = \mathbf{D}_{d,correct} + \mathbf{R}_{no}, \text{ and}$$

$$\mathbf{T}_{d,incorrect} = \mathbf{D}_{d,incorrect} + \mathbf{R}_{yes}.$$

The expected values are,

$$\mu_{d,correct} = E[\mathbf{D}_{d,correct}] + E[\mathbf{R}_{no}],$$

$$\mu_{d,incorrect} = E[\mathbf{D}_{d,incorrect}] + E[\mathbf{R}_{yes}].$$

For the **two-distractor condition** dd , there are three cases. The first case is when both distractors are processed correctly, in parallel. For this case, the response time is determined by the last component process completed. The probability that Case 1 occurs is,

$$p_{dd,Case1} = P(\mathbf{Z}_d = 1 \text{ and } \mathbf{Z}_d = 1) = p_d^2.$$

The correct processing time for Case 1 is,

$$T_{dd,Case1} = \max \{D_{d_1,correct}, D_{d_2,correct}\} + R_{no},$$

with

$$\mu_{dd,Case1} = E[\max \{D_{d_1,correct}, D_{d_2,correct}\}] + E[R_{no}].$$

The second case is when one distractor is processed incorrectly and the other distractor is processed correctly, in which case the processing is terminated when the distractor is processed incorrectly. The probability that Case 2 occurs is,

$$p_{dd,Case2} = P(Z_{d_1} = 0 \text{ and } Z_{d_2} = 1 \text{ or } Z_{d_1} = 1 \text{ and } Z_{d_2} = 0) = 2p_d(1 - p_d).$$

The incorrect processing time for Case 2 is,

$$T_{dd,Case2} = D_{d,incorrect} + R_{yes},$$

with

$$\mu_{dd,Case2} = E[D_{d,incorrect}] + E[R_{yes}].$$

The third case is when both distractors are processed incorrectly. In this case, the response time is determined by the fastest of the two component processes. The probability that Case 3 occurs is,

$$p_{dd,Case3} = P(Z_{d_1} = 0 \text{ and } Z_{d_2} = 0) = (1 - p_d)^2.$$

The incorrect processing time for Case 3 is,

$$T_{dd,Case3} = \min \{D_{d_1,incorrect}, D_{d_2,incorrect}\} + R_{yes},$$

with

$$\mu_{dd,Case3} = E[\min \{D_{d_1,incorrect}, D_{d_2,incorrect}\}] + E[R_{yes}].$$

Only the first case yields a correct response, thus

$$p_{dd} = p_{dd,Case1} = p_d^2, \text{ and}$$

$$T_{dd,Case1} = \max \{D_{d1,correct}, D_{d2,correct}\} + R_{no},$$

with

$$\mu_{dd,Case1} = E[\max \{D_{d1,correct}, D_{d2,correct}\}] + E[R_{no}].$$

The other two cases yield incorrect responses, yielding a mixture distribution. The weight for Case 2 is the fraction of incorrect responses due to Case 2 relative to all incorrect responses,

$$w_{Case2} = \frac{2p_d(1-p_d)}{(2p_d(1-p_d) + (1-p_d)^2)} = \frac{2p_d}{1+p_d}.$$

Similarly, the fraction of incorrect responses due to Case 3 relative to all incorrect responses is,

$$w_{Case3} = \frac{(1-p_d)^2}{(2p_d(1-p_d) + (1-p_d)^2)} = \frac{1-p_d}{1+p_d}.$$

The weighted combination of Cases 2 and 3 yields the expected incorrect response time,

$$\begin{aligned} \mu_{dd,incorrect} &= w_{Case2}E[T_{dd,Case2}] + w_{Case3}E[T_{dd,Case3}] \\ &= \left(\frac{2p_d}{1+p_d}\right)(E[D_{d,incorrect}] + E[R_{yes}]) \\ &\quad + \left(\frac{1-p_d}{1+p_d}\right)(E[\min \{D_{d1,incorrect}, D_{d2,incorrect}\}] + E[R_{yes}]) \\ &= \left(\frac{2p_d}{1+p_d}\right)(E[D_{d,incorrect}]) + \left(\frac{1-p_d}{1+p_d}\right)(E[\min \{D_{d1,incorrect}, D_{d2,incorrect}\}]) + E[R_{yes}]. \end{aligned}$$

For the **one target and one distractor condition** td , there are four mutually exclusive cases. The cases are distinguished by whether the target is processed correctly or incorrectly, coupled with whether the distractor is processed correctly or incorrectly.

The first case is that the target is processed correctly and the distractor is processed correctly. Here, only the processing of the target determines the response time. The probability that Case 1 occurs is,

$$p_{td,Case1} = p_t p_d$$

$$T_{td,Case1} = D_{t,correct} + R_{yes},$$

with

$$\mu_{td,Case1} = E[D_{t,correct}] + E[R_{yes}].$$

The second case is that the target is processed correctly and the distractor is processed incorrectly. Now, there is a race between the two processing times. The probability that Case 2 occurs is,

$$p_{td,Case2} = p_t (1 - p_d),$$

$$T_{td,Case2} = \min \{D_{t,correct}, D_{d,incorrect}\} + R_{yes},$$

with

$$\mu_{td,Case2} = E[\min \{D_{t,correct}, D_{d,incorrect}\}] + E[R_{yes}].$$

The third case is that the target is processed incorrectly and the distractor is processed correctly. Now both processes must complete to determine a response. The probability that Case 3 occurs is,

$$p_{td,Case3} = (1 - p_t) p_d,$$

$$T_{td,Case3} = \max \{D_{t,incorrect}, D_{d,correct}\} + R_{no},$$

with

$$\mu_{td,Case3} = E[\max \{D_{t,incorrect}, D_{d,correct}\}] + E[R_{no}].$$

The fourth case is that the target is processed incorrectly, and the distractor is processed incorrectly. Here, the distractor processing determines the response time. The probability that Case 4 occurs is,

$$p_{td,Case4} = (1 - p_t)(1 - p_d),$$

$$T_{td,Case4} = D_{d,incorrect} + R_{yes},$$

with

$$\mu_{td,Case4} = E[D_{d,incorrect}] + E[R_{yes}].$$

The probability of a correct response is achieved through the mutually exclusive cases 1, 2, and 4. It is the probability that the target is processed correctly, as in Case 1 or Case 2, plus the probability that the distractor is processed incorrectly as in Case 4,

$$\begin{aligned} p_{td} &= p_t p_d + p_t(1 - p_d) + (1 - p_t)(1 - p_d) \\ &= 1 - p_d(1 - p_t). \end{aligned}$$

For the three cases that contribute to a correct response, the weights are

$$\begin{aligned} w_1 &= \frac{p_{td,Case1}}{p_{td}} = \frac{p_t p_d}{(1 - p_d(1 - p_t))}, \\ w_2 &= \frac{p_{td,Case2}}{p_{td}} = \frac{p_t(1 - p_d)}{(1 - p_d(1 - p_t))}, \\ w_4 &= \frac{p_{td,Case4}}{p_{td}} = \frac{(1 - p_t)(1 - p_d)}{(1 - p_d(1 - p_t))}. \end{aligned}$$

Using these weights, the expected correct response time is,

$$\mu_{td,correct} = w_1 E[T_{td,Case1}] + w_2 E[T_{td,Case2}] + w_4 E[T_{td,Case4}]$$

which does not simplify nicely.

The expected incorrect response time is solely due to Case 3, and is,

$$\mu_{td,Case3} = E[\max \{D_{t,incorrect}, D_{d,correct}\}] + E[R_{no}].$$

The standard self-terminating, unlimited-capacity, parallel model with errors does not allow quantitative predictions, as in the corresponding serial model. The minimum and maximum terms prevent such specific quantitative predictions. Nevertheless, one can make the qualitative prediction that the response time is always reduced with two targets compared to one.

Standard Self-terminating, Fixed-capacity, Parallel Model with Errors

We next consider a fixed-capacity version of our standard self-terminating, parallel model with errors. The term *fixed capacity* is from information theory (Taylor et al., 1967) and means that a constant amount of information is processed from the entire set of stimuli. Thus, assuming equal allocation, half as much information can be extracted from each of two stimuli as can be extracted from one stimulus alone. This idea can be implemented using a sampling process described by (Shaw, 1980). Such a fixed-capacity model is a special case of a limited-capacity model. The fixed-capacity parallel model provides a useful landmark among the wide range of possible limited-capacity models (White et al., 2018).

Our goal in analyzing this model is to determine if it predicts that redundant targets have faster response time than single targets, such as found for our unlimited-capacity parallel model. Alternatively, capacity limits might overwhelm the

redundancy gain and result in slower response times as found for our serial model. Establishing this prediction helps distinguish serial and parallel models in general.

One can start from the parallel model described in the preceding section by replacing the unlimited-capacity assumption with fixed capacity. Unfortunately, we know of no way to analyze such a distribution-free model. Instead, we define two special cases of the fixed-capacity model and derive numerical predictions. To foreshadow the results, the two versions have quite different predictions.

Version 1: Simple diffusion processes.

To begin, assume all of the model structure of the previous unlimited-capacity parallel model and add a specific stochastic process that generates the responses and response times. In this version, we use a simple diffusion process that has often been applied to response time (Palmer et al., 2005) and has been elaborated as a theory of visual search (Corbett & Smith, 2020).

Consider n stimuli, where each stimulus can be a target t or a distractor d . For the current task, a diffusion process applied to response time describes the continuous accumulation of relative evidence for the presence of the target t versus the presence of a distractor d . Let the accumulated evidence for stimulus $i, i = 1, \dots, n$, correspond to a random variable that varies over time $\mathbf{U}_i(x)$ where x is time. At time zero, $\mathbf{U}_i(0) = 0$. As time increases, evidence is accumulated from a target at a mean rate r_t and from a distractor at a mean rate $-r_d$. (For the details of representing evidence as a signal-to-

noise ratio, see Palmer, et al., 2005). The change for the fixed-capacity model is that the rate for this model is reduced by a factor of $1/\sqrt{2}$ relative to an unlimited-capacity model. This is the result of the rate being determined by a set of independent samples that are equally allocated when there are multiple stimuli (Shaw, 1980). With two stimuli, half as many samples can be allocated, as to a single stimulus. This results in twice the variability, or $\sqrt{2}$ the standard deviation of the estimate of the stimulus information. This scales the effect of the stimulus by $1/\sqrt{2}$. For example, if r_t and r_d are 2.0 for the unlimited capacity model, they would be 1.41 for the fixed-capacity model.

The response occurs by evaluating the accumulated evidence for each stimulus. Starting at time zero, all stimuli are unlabeled, and as time increases, the evidence is evaluated to label the stimuli with a positive or negative decision, as follows.

- At time step x , for unlabelled stimuli, evaluate $U_i(x)$:
 - If $U_i(x) > a$, then label stimulus i with a positive decision, and terminate with a “yes” response.
 - If $U_i(x) < -b$, then label stimulus i with a negative decision, and continue.
- If all stimuli are labeled with a negative decision, terminate with a “no” response.

Otherwise, increment the time step and repeat.

To complete the definition of the fixed-capacity diffusion model, we fix the coefficient of variability of the residual time to 0.1. This value for the coefficient of variability is motivated by the idea that the residual processes are stereotyped and have

a relatively low variance. In contrast, the component processing time from the diffusion process typically has a much larger coefficient of variability of around 0.8 to 1.0. The variability of the total response time is the sum of the variability of the residual time and the component processing time. The coefficient of variability for the total response time found in perceptual tasks varies from 0.1 for strong stimuli (dominated by the residual time) to 0.5 for weak stimuli (contributions from both residual time and component processing time). While useful for specifying the models, this residual time parameter has no effect on the redundant target effect.

To calculate the predictions of this model, we first choose parameter values relative to Experiment 2 semantic condition. To do this, the experiment is summarized by the mean percent correct, mean correct response time, and standard deviation for the correct response times for the single target and single distractor conditions. This gives six statistics describing the data that are listed in Table A1. In this table, the values for accuracy and mean response time have been reported in the body of the paper. Here we also add the value for the standard deviation of response time calculated as the mean of the standard deviation for each participant. It is sometimes useful to restate these values in terms of the coefficient of variation (standard deviation/mean). For targets, the coefficient of variation was 0.26 for Experiment 2 and was 0.30 for Experiment 3. These values are typical for response time tasks with relatively difficult

discriminations (Luce, 1986, p 64-66, and p 208-211). These estimates of variability are important to constrain specific models of response time.

Table A1

Properties used to determine model parameters

	Experiment 2	Experiment 3
<u>Single Targets</u>		
Percent correct	90.8%	97.8%
Mean response time	665 ms	641 ms
Standard deviation of response time	172 ms	190 ms
<u>Single Distractors</u>		
Percent correct	93.9%	82.6%
Mean response time	693 ms	806 ms
Standard deviation of response time	183 ms	207 ms

The six model parameters are estimated from the six statistics. The estimated parameter values are listed in Table A2. Importantly, all of this is done with just the single stimulus condition in Experiment 2. Finally, using these parameters, we calculate the predicted effects of two targets compared to a single target. For this experiment,

this fixed-capacity model predicts that two targets are faster and more accurate than a single target (gain of 28 ms and 6.2% correct).

Table A2

Parameters for the Diffusion Model

Parameter	Symbol	Experiment 2 Values	Experiment 3 Values
rate for a target	r_t	2.33	1.92
rate for a distractor	r_d	1.98	2.24
upper bound	a	0.672	0.338
lower bound	$-b$	0.504	0.931
mean residual time for a “yes” response	$E[R_{yes}]$	0.418	0.458
mean residual time for a “no” response	$E[R_{no}]$	0.478	0.461

Next consider parameters based on the data from Experiment 3 semantic condition. As above, we use six statistics from Table A1 to determine the six parameter values. The estimated parameters are given in Table A2. For these conditions, this

fixed-capacity model also predicts that two targets are faster and more accurate than one (gain of 59 ms and 1.9% correct). Thus, these conditions also yield positive redundant target effects. Indeed, for all conditions that avoid extreme parameters, the redundant target effect remains positive for this version of the model.

We also examined the effect of redundant targets on the standard deviation of response time. For Experiment 2, the predicted standard deviation was smaller for two targets (0.161 s) compared to one target (0.172 s). For Experiment 3, the predicted standard deviation was smaller for two targets (0.134 s) compared to one target (0.190 s). Thus, for this version of the fixed-capacity model there was no sign of a larger standard deviation for the two-target condition compared to the one target condition.

In summary, we evaluated a version of the fixed-capacity parallel model that depends on a simple diffusion process. For all conditions expected in a typical experiment, the fixed-capacity parallel model predicts a positive effect of redundant targets. Thus, the predictions of this version of a fixed-capacity parallel model are distinct from the predictions of negative redundant target effects made by our standard self-terminating serial model.

Version 2: Linear ballistic accumulators.

A different model of response time is the linear ballistic accumulator (LBA) model (Brown and Heathcote, 2008). It differs from the simple diffusion model in several ways. First, the stochastic element is variability in the rate of accumulation from

trial to trial, instead of from moment to moment. Second, there are separate accumulators for each response, instead of comparing positive and negative evidence within a single accumulator. Third, there is a “bias” contribution to the rate of accumulating evidence for each accumulator, instead of separate bounds for the net positive and negative evidence. We implemented a particularly simple version of this model. The variability of the rate parameter was described by a Gamma distribution (Terry et al., 2015) to avoid the complications of negative rates that occurred in the original formulation arising from using a Gaussian distribution. In addition, we dropped variability in the start point. Finally, the predictions from the redundant target conditions were implemented following the derivation in (Eidels et al., 2010).

For this model, the accumulated evidence for a “yes” or “no” response is in separate accumulators, denoted Y or N , respectively. Such a pair of accumulators exists for each stimulus $i, i = 1, \dots, n$. The accumulated evidence for stimulus i corresponds to two random variables, denoted $\mathbf{U}_{Y_i}(x)$ and $\mathbf{U}_{N_i}(x)$, where x is time. At time zero, $\mathbf{U}_{Y_i}(0) = 0$ and $\mathbf{U}_{N_i}(0) = 0$. Evidence is accumulated for a target at a rate of r_{t_Y} and r_{t_N} for each Y and N accumulator, respectively. Similarly, evidence is accumulated for a distractor at a rate of r_{d_Y} and r_{d_N} for each Y and N accumulator, respectively. These rates can be interpreted in terms of signal component and bias component. Let the signal for a target be $r_{t,signal} = r_{t_Y} - r_{t_N}$ and the bias for a target be $r_{t,bias} = r_{t_N}$.

Similarly, let the signal for a distractor be $r_{d,signal} = r_{d_N} - r_{d_Y}$ and the bias for a distractor be $r_{d,bias} = r_{d_Y}$.

This separation of signal and bias components is needed to introduce the idea of fixed capacity. The signal rates determine the accuracy of the response. For example, if $r_{t,signal} = 0$, the responses on the target trials are at chance. To incorporate fixed capacity, the signal rates are reduced by a factor of $1/\sqrt{2}$, just as was done in the diffusion model.

Two additional parameters for this model are: a common bound for all accumulators b , and a common standard deviation for the variability of all rate parameters r_{SD} . These two parameters and all of the rate parameters share a common factor, so one can fix one of these parameters. Hence, we set $r_{SD} = 1$, which is equivalent to making all of these parameters relative to the standard deviation of the rate (see Palmer, et al., 2005). In addition, to further reduce the number of parameters, we set the bound $b = 1$. This is possible because the $r_{t,bias}$ and $r_{d,bias}$ parameters act in a similar way to having separate bound for both accumulators.

A response occurs when the evidence in any “yes” accumulator $\mathbf{U}_{Y_i}(x)$ reaches the bound for a “yes” response, or all of the “no” accumulators reach the bound for a “no” response. Starting at time zero, all stimuli are unlabeled, and as time increases, the evidence is evaluated to label each stimulus with a positive or negative decision, as follows.

- At time step x , for unlabelled stimuli, evaluate $\mathbf{U}_{Y_i}(x)$: and $\mathbf{U}_{N_i}(x)$:
 - If $\mathbf{U}_{Y_i}(x) > b$, then label stimulus i with a positive decision, and terminate with a “yes” response.
 - If $\mathbf{U}_{N_i}(x) > b$, then label stimulus i with a negative decision, and continue.
- If all stimuli are labeled with a negative decision, terminate with a “no” response. Otherwise, increment the time step and repeat.

Finally, to complete the model, we add two residual time parameters: the mean of the residual time for a “yes” response $E[\mathbf{R}_{yes}]$ and the mean of the residual time for a “no” response $E[\mathbf{R}_{no}]$. Together, there are six parameters and they are listed in Table A3.

Table A3

Parameters for the LBA Model

Parameter	Symbol	Experiment 2 Values	Experiment 3 Values
signal rate for a target	$r_{t,signal}$	2.11	1.45
signal rate for a distractor	$r_{d,signal}$	1.76	2.23
bias rate for a target	$r_{t,bias}$	0.593	1.23
bias rate for a distractor	$r_{d,bias}$	0.925	0.187
mean residual time for a “yes” response	$E[R_{yes}]$	0.251	0.209
mean residual time for a “no” response	$E[R_{no}]$	0.269	0.351

As with the diffusion model, we numerically solve for these parameters based on six statistics from the single target and single distractor conditions of our experiments given in Table A1. The values of these six parameters are listed in Table A3. From

these parameters, the predicted redundant target effects were calculated. For Experiment 2 semantic condition, there was a negative redundant target effect of -1 ms on response time, and a positive 8.0% effect on accuracy. For Experiment 3 semantic condition, there was a negative redundant target effect of -40 ms on response time, and a positive 2.1% effect on accuracy. The important new result is that this model can predict a negative redundant target effect.

Regarding the standard deviation of the response times, there are similar results. For Experiment 2, the predicted standard deviation is smaller for two targets (0.162 s) compared to one target (0.172 s). For Experiment 3, the predicted standard deviation is larger for two targets (0.213 s) compared to one target (0.190 s). Thus, this model can predict both possible effects on the standard deviation of response time.

Why do these two versions of the fixed-capacity, parallel model differ regarding the sign of the redundant target effect? We suspect an important factor is the degree of variability in the component processing time for each stimulus. For the diffusion model, this variability is quite high, with a coefficient of variability of the component time around 1.0. In contrast, for the linear ballistic accumulator model, the variability is determined by the set of parameters. For Experiment 3, the coefficient of variability of the component time for the selected parameters was only 0.25. Such lower variability results in a smaller redundant target effect that can be overcome by the increase in processing time due to fixed capacity. This results in a negative redundant target effect.

Summary

We have presented two version of the standard self-terminating, fixed-capacity, parallel model with errors. One is based on a diffusion process and the other based on the linear ballistic accumulator. For parameters based on the single stimulus conditions of Experiment 3 with the semantic task, these models make quite different predictions about the redundant target effect on response time: one positive and the other negative. Thus, this general class of model makes no prediction about whether the redundant target effect for response time is positive or negative.

Details about Calculating Predictions of Redundant Target Effects

In the introduction of the article, Figure 2 shows typical range of predictions of redundant target effects on mean correct response time for our three landmark models. Here, we describe the details of how we made these predictions.

The predictions for the standard serial model are relatively simple to calculate, based on Equation (7). The only factors affecting the predictions are the mean component processing time for errors on target trials (misses), and the proportion of correct responses for a single target. For all predictions and all models, we kept the mean correct response time fixed at 600 ms and percent errors at 5%. To set the mean component processing time for errors, we also assumed a mean residual time of 200 ms. Then, for the upper end of the range, we assumed equal component processing time for correct (hit) and error (misses) responses. For the lower end of the range, the mean

component processing time for errors was assumed to be twice as long as the mean component processing time for correct responses. These predictions span the range expected in a typical experiment.

For our standard unlimited-capacity, parallel model, there are no direct numerical predictions from Equation (17). Instead, we rely on the two special cases we have already developed: the diffusion model and the linear ballistic accumulator model. For the upper end of the range, we used the diffusion model with parameters that for the single stimulus trials yield 5% errors, a mean response time of 600 ms and coefficient of variability of 0.3. These constraints were used for both target and distractor trials. The result is a relatively large redundant target effect. For the lower end of the range, we used the linear ballistic accumulator model in a similar way. To minimize the redundant target effects, we reduced the coefficient of variability to 0.1. This yielded a relatively small redundant target effect. While these predictions are not bounds, they illustrate the range of typical effects.

Lastly, consider the predictions of our standard fixed-capacity, parallel model. As with our unlimited-capacity model, we rely on the two special cases with similar constraints. The one twist is that at the lower end of the range, the linear ballistic accumulator model was paired with a coefficient of variability of 0.3. For this model, the predictions now span a range that includes both positive and negative redundant target effects.

Predictions of redundant target effects on response time distributions

As explained in preceding sections of this Appendix, the serial model makes unique predictions about the distributions of correct response times (RTs) on redundant target trials. That is because on some fraction of trials, the target that gets processed first is misidentified as a distractor, so search continues to process the second target, which is then responded to correctly with a relatively long latency. In theory, this could create a bimodal distribution of RTs, with a greater standard deviation than the distribution of RTs on single-target trials. But would such a change in the RT distribution be measurable in our experiments?

To illustrate, we conducted simulations of RT distributions based on the semantic task of Experiment 3. For this simulation, we increased the probability of an error in the one-target condition from the empirical mean of 2.2% to 10%. that increases the distinctiveness of the prediction of the standard serial model, which depends on initial misidentifications of one target before processing the other.

The following **Figure A1** shows four distributions. We start with a distribution predicted by the linear ballistic accumulator (LBA) model fit to the one-target condition. This distribution for the one-target condition is the solid black curve in the figure. It matches the data in both mean and standard deviation, but its exact shape is from the LBA model.

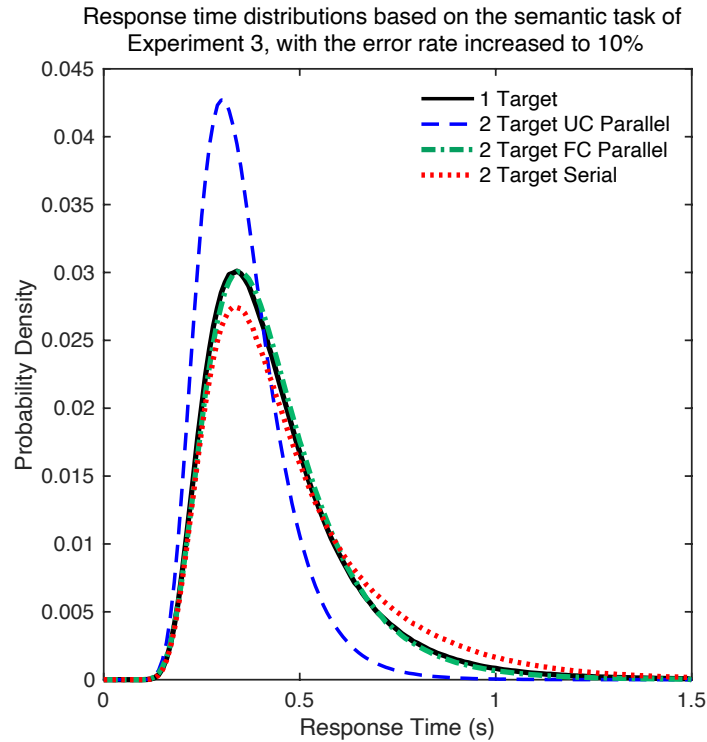


Figure A1: Predicted distributions of correct response times for trials with 1 target (black) or for trials with 2 targets, for 3 different models in different colors. UC = unlimited capacity; FC = fixed capacity.

Given this starting point, the three standard models predict specific distributions for the two-target condition. The prediction for the unlimited-capacity parallel model is the dashed blue curve. It shows the usual positive redundant target effect: a shift to faster response times. Notice the reduction in the tail of the distribution. The prediction for the fixed-capacity parallel model is the dot-dashed green curve. It shows a tiny negative redundant target effect but, for this example, mostly falls near the one-target distribution. The prediction for the standard serial model is the dotted red curve. It shows a somewhat larger shift to slower response times. Its unique feature is the tail of the distribution which is heavier on the right than the other models' predictions.

Unfortunately, even with the error rate exaggerated at 10%, the serial model's prediction of a heavier tail is a small effect. It is nothing like a response time distribution with two modes. This is because the variability of response time is large relative to the duration of the additional processing time in the redundant target condition. Detecting the heavier tail predicted by the standard serial model would require a large and careful experiment and is beyond the current experiments.

Rather than comparing the shapes of the response time distributions, a simpler approach is to investigate the variability of the distributions. In the preceding Appendix sections that describe each model separately, we attempted to derive predictions for the standard deviation (SD) of responses times. The standard serial model predicts that the SD in the redundant target condition is larger than in the single target condition. The fixed-capacity parallel model can predict either a larger or smaller SD in the redundant target condition. For the unlimited capacity parallel model, we have not been able to formally derive a prediction. But for all the special cases we have investigated, the SD in the redundant target condition is smaller than in the single target condition.

In summary, the RT distributions could in theory confirm the results in the RT means. The current experiments were not designed to reveal these effects, which our simulations predict would be small relative to the observed variability within each condition. Moreover, the predictions for RT variability exactly mirror the predictions for

the means, for all three classes of models. Therefore, they would not add any unique leverage to distinguish the models.